

Stellar Structure & Evolution

Charlie Conroy

v1, September 6, 2026

Preface

This book forms the core of AY204, a graduate course I have taught at Harvard since 2018. Principal sources for this material include “Understanding Stellar Evolution” by Lamers & Levesque, “Stellar Interiors” by Hansen, Kawaler, and Trimble, “Principles of Stellar Evolution and Nucleosynthesis” by Clayton, the lecture notes from Onno Pols, Lars Bildsten, Rich Townsend, and the Princeton graduate course on stellar evolution. I thank the former students of AY204 for providing feedback and identifying errors in previous versions.

Contents

1	Overview	1
2	Radiation & Matter	8
3	Stellar Atmospheres	15
4	Equation of State	25
5	Opacity & Nuclear Energy Generation	34
6	Mechanical Structure & Energy Conservation	44
7	Energy Transport	52
8	Simple Stellar Models	58
9	Fully Convective Stars	64
10	Post-Main Sequence Evolution	70
11	Stellar Winds & Mass Loss	78
12	Evolution of Low-Mass Stars	83
13	Evolution of Massive Stars	91
14	Binary Stars I	97
15	Binary Stars II	102
16	Stellar Variability	109
17	Asteroseismology	117
18	White Dwarfs	126

19 Neutron Stars & Black Holes	134
20 Supernovae I	140
21 Supernovae II	146
22 Gravitational Waves	153

Chapter 1

Overview

Stars are the basic building blocks of the universe and are responsible for the production of most elements via nucleosynthesis. One of the greatest achievements of 20th century science was the development of a detailed, quantitative theory for stellar structure and evolution. This course will cover both this remarkable achievement and the key outstanding questions in this field. An emphasis will be placed on fundamental concepts and connections to observations.

This lecture will review the basic properties of stars and how they are measured. We will also review the key timescales relevant for stellar evolution. You would do well to thoroughly understand the timescales, as we will return to them throughout the course.

Let's begin by describing the basic parameters of stars and how they are measured.

- Mass (M). This is by far the most fundamental parameter of a star, largely setting its evolution, appearance, etc. The canonical range of stellar masses is $0.08 < M < 150 M_{\odot}$, where the lower limit is the hydrogen-burning limit, and the upper limit is not well-determined from either theory or observations (literature values range from $100 - 300 M_{\odot}$ or the upper limit). Unfortunately, stellar masses are very difficult to measure. The most reliable constraints require dynamics. Eclipsing binaries offer the most powerful constraints (masses measured to a precision of $\lesssim 3\%$), but such configurations are rare (of order a few hundred are known). Masses for much larger samples of stars can be estimated by appealing to stellar models, as we will see later in this course.
- Age (t). This parameter tells us where along its evolutionary history a star of a given

mass is going to be. The parameter can range from ≈ 0 to the age of the Universe (≈ 14 Gyr). Like mass, age is also very difficult to measure. The only truly fundamental technique is via radioactive decay of long-lived isotopes such as U and Th. This technique has been applied to very few stars. The most common technique is the use of stellar models, although there are also empirical techniques such as stellar spindown (gyrochronology) and the Lithium depletion boundary. An excellent review of the measurement of stellar ages is provided in Soderblom (2010).

- **Metallicity (Z).** The metal-content of a star has an important effect on both the internal structure and observable appearance of stars due to several effects including opacity and a change to the mean molecular weight (we'll get to these concepts later). Z is defined as the mass fraction in metals (elements heavier than helium), while X and Y are the mass fraction in H and He. By definition $X + Y + Z = 1$. The surface metallicity of a star can be determined via spectroscopy, but the relation between surface and bulk metallicity is not always so clear (and requires models!). In principle, measuring Z is a tall order, as it requires measuring the abundances of *all* elements in the star. In practice, one often measures just a few elements (e.g., Fe, O) and makes some reasonable assumptions regarding the behavior of the other elements.

Notation: you will see many different ways of writing metallicity, including Z , $Z/Z_{\odot}$, and $[Z/H] \equiv \log (N_Z/N_H) - \log (N_Z/N_H)_{\odot}$, where $\odot$ refers to the Sun. The latter is close to, but not identical to $\log Z/Z_{\odot}$ (though they are often used interchangeably). You will also encounter generalizations of this notation to: $[A/B] \equiv \log (N_A/N_B) - \log (N_A/N_B)_{\odot}$. In words, $[A/B]$ is the logarithm of the ratio of the number of atoms of elements A and B relative to the solar composition. This might sound almost comically opaque, but we will see later in the course the value of this notation.

- **Composition.** Metallicity above is in fact too crude - stellar evolution depends not only on Z but also on the detailed mixture of different kinds of elements. This is sometimes called the composition, or abundance pattern of a star.

In terms of computing stellar models, the above list are the main parameters determining everything else. We have ignored at least three other relevant variables including rotation, binarity, and magnetic fields. We will return to these topics in the second half of the course.

Question: Try to come up with some other ways that we could measure the masses or ages of stars via a relatively fundamental technique (i.e., without appealing to stellar models).

There are of course many more useful quantities that can be measured from a star. You should think of these quantities as “derived” in the sense that they are predicted once the first four parameters listed above are known. Most of the parameters listed below are much easier to measure than the ones listed above.

- Bolometric luminosity (L_{bol}). “Bolometric” means total, i.e., integrated over all frequencies. Stars lie in the range $10^{-3} < L_{\text{bol}} < 10^6 L_{\odot}$. This requires measuring the flux over all frequencies (or wavelengths), which is generally difficult. Thankfully most stars emit most of their radiation in the UV-NIR. Luminosities also require a measure of the distance to the star, which has historically been a major industry, and was revolutionized in 2018 with the release of high-quality parallaxes for $> 1\text{B}$ stars from the ESA *Gaia* satellite.
- Stellar radius (R). Ranges from $\sim 0.1 < R < 10^3 R_{\odot}$. Radii are difficult to measure, and can be done either with eclipsing binaries or via interferometry. The latter is limited to very nearby sources.
- Surface gravity ($\log g$). Stellar spectra offer a relatively direct measurement of $\log g$, and it is a natural variable when considering plane-parallel atmospheres, which is a good approximation for many stars. Recall $g \equiv \frac{GM}{R^2}$. Ranges from $-1 < \log g < 5$. For historical reasons $\log g$ is quoted without units but cgs units are assumed.
- Effective temperature (T_{eff}). The temperature of stars ranges from approximately $2000 < T_{\text{eff}} < 50,000$ K. T_{eff} can be determined from spectra or just as easily from broadband colors of stars. The effective temperature is *defined* via the Stefan-Boltzmann Law: $L = 4\pi R^2 \sigma_{\text{SB}} T_{\text{eff}}^4$, where σ_{SB} is a constant.

A large and ever-growing industry is devoted to the measurement of these quantities.

Figure 1.1 shows an Hertzsprung-Russell (HR) diagram (T_{eff} vs. $\log L_{\text{bol}}$). I’ve noted a few important features including the zero-age main sequence (ZAMS) and evolutionary paths for a solar-type star and a massive ($M = 40M_{\odot}$) star. Important phases include the red giant branch (RGB), red clump (RC), tip of the RGB (TRGB), and the white dwarf cooling sequence.

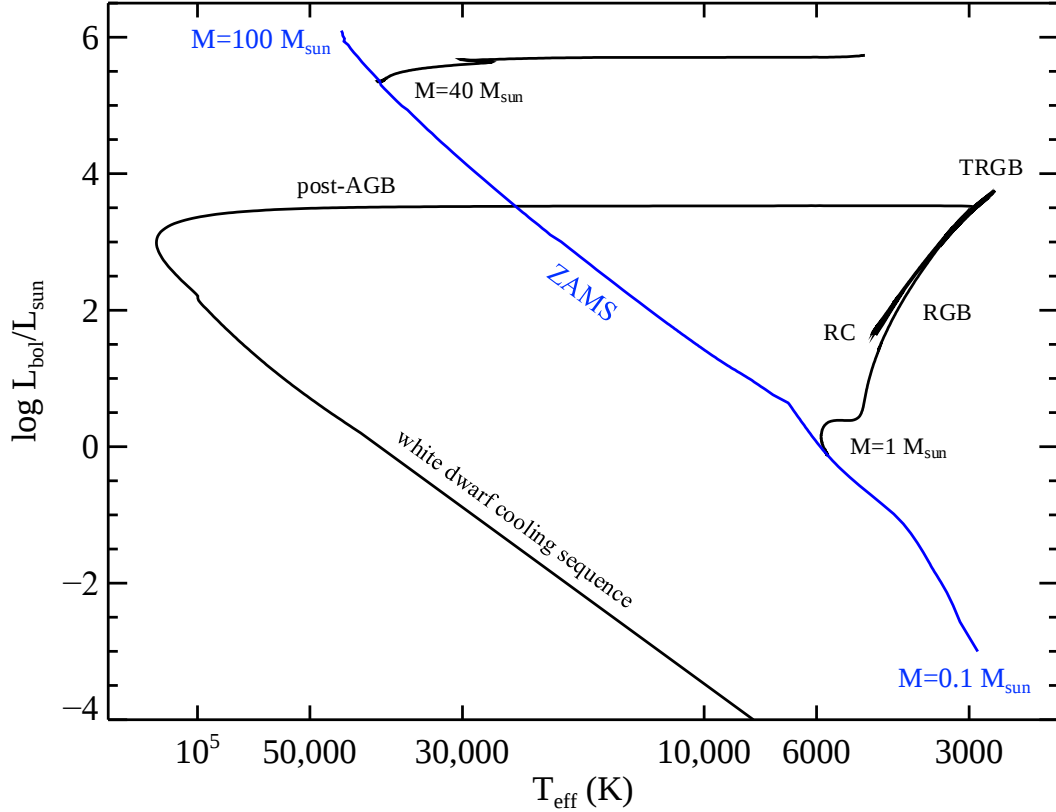


Figure 1.1: HR diagram for stars at solar metallicity. The zero-age main sequence (ZAMS) is shown for stars ranging in mass from $0.1 - 100M_{\odot}$. Evolutionary tracks are also shown for $1M_{\odot}$ and $40M_{\odot}$ stars. Several important evolutionary phases are labeled for the solar mass model.

As we will see, the evolution of stars above and below $\approx 8 - 10 M_{\odot}$ is very different. There are other important transition points at $\approx 0.3 M_{\odot}$ and $\approx 2 M_{\odot}$, which you can identify by the changes in the slope of the ZAMS curve.

The observational analog of the HR diagram is the color-magnitude diagram (CMD), in which a color is plotted on the x-axis, and absolute magnitude is plotted on the y-axis. But what is a magnitude?

We define an astronomical magnitude m via:

$$m \propto -2.5 \log f \sim -\ln f \quad (1.1)$$

where f is the flux. It is important to understand that there is nothing profound about the number -2.5 ; the number arises from an attempt to place the historical naked-eye

observations (in which stars were given numerical rankings by brightness) on a firmer footing. Our eyes are approximately log-sensitive to flux, which motivates why a linear change in magnitude corresponds to a logarithmic change in flux.

The above equation requires a constant, also known as a zero-point. There are two common conventions in use today: 1) use the star Vega to define the zero-points such that $m \equiv 0$ for all wavelengths; 2) the AB system, in which $m = -2.5 \log f_\nu - 48.60$ (Oke & Gunn 1983). The AB system has the convenient property that a magnitude corresponds to a physical unit of flux. The numerical value of the zero-point is arbitrary and was chosen to correspond to the Vega system for the V -band filter.

The equation above holds both for in a monochromatic sense, e.g., $m_\nu \propto -2.5 \log f_\nu$ and also if one takes observations through a finite bandpass filter (literally a piece of glass placed somewhere in the optical path of the instrument). For simplicity, assume that a bandpass filter is box-shaped with a central wavelength, λ_{eff} and a width $R \equiv \lambda_{\text{eff}}/\Delta\lambda$. For many astronomical filters $R \sim 5$. There are many filter systems in use, including Johnson-Cousins (UBVRI), SDSS (ugriz), and NIR systems (JHKLM). There are hundreds of filters spanning the UV-IR.

We can use the magnitude equation to obtain a relation between magnitudes and distances:

$$m_1 - m_2 = -2.5 \log (f_1/f_2) = 5 \log (d_1/d_2) \quad (1.2)$$

The absolute magnitude, M , is defined as the magnitude a star would have it were placed at 10 parsecs, so we have:

$$m - M = 5 \log (d/10 \text{ pc}) \quad (1.3)$$

where $m - M$ is known as the distance modulus (sometimes referred to as μ).

Question: What are advantages and disadvantages of the Vega and AB zero-point systems? How do we measure the absolute flux of anything in the sky? How accurately do you think we know the absolute flux of the best-studied sources?

There are several important timescales relevant for stellar evolution associated with the changes in the mechanical structure, thermal structure, and changes in composition.

The **dynamical timescale** is given by:

$$t_{\text{dyn}} = \sqrt{\frac{R^3}{GM}} \quad (1.4)$$

which for the Sun is ≈ 1600 s. The dynamical time can be thought of as the time it would take for a sound wave to travel across the star. It is the timescale on which a star can react to perturbations from hydrostatic equilibrium. The dynamical time is similar to the free-fall time, which is the time it would take for a star to collapse under gravity if the sources of support were suddenly removed.

The **thermal timescale** (also known as the Kelvin-Helmholtz timescale) is the time required for a star to radiate its current gravitational binding energy:

$$t_{\text{KH}} = \frac{E}{\dot{E}} = \frac{E}{L} = \frac{G M^2 / R}{L} \quad (1.5)$$

which for the Sun is $\approx 3 \times 10^7$ yr. This timescale played an important role in understanding the source of energy production in the Sun. In the 1860s Lord Kelvin argued based on the above timescale that the age of the Sun (and hence the Earth) was $< 10^8$ yr. There was a fascinating debate around this time related to the fidelity of geological dating (a field in its infancy at the time). Once it was firmly established that the Earth was older than 10^8 yr it was clear that the Sun must have an energy source other than gravitational contraction.

The **nuclear timescale** is the time it takes for a star to burn through its nuclear fuel:

$$t_{\text{nuc}} = \frac{E_{\text{nuc}}}{L} = \frac{f 0.007 M c^2}{L} \quad (1.6)$$

where 0.007 is the fraction of the rest mass of H converted to energy as hydrogen is converted to helium, and f is the fraction of the total star that is available for nuclear burning, which is ≈ 0.1 for the Sun. This leads to $t_{\text{nuc}} \approx 10^{10}$ yr for the Sun.

A final important timescale is the **photon-diffusion timescale**. This is the time it takes for a photon to travel from the center to the surface of a star via diffusion. Diffusion is a random-walk process, with N steps each of length λ . The key feature of a random walk is that the distance traveled after N steps is $d_N = \sqrt{N}\lambda$. If photon collisions with matter are mediated by electron (Thomson) scattering then the mean collision time is $t_{e\gamma} = \bar{\lambda}/c$ where $\bar{\lambda} = (n_e \sigma)^{-1}$ and σ is the Thomson cross section. For the Sun, $\bar{\lambda} \approx 1$ cm (assuming $\bar{\rho} = \frac{3M}{4\pi R^3}$ and pure hydrogen composition). Setting $d_N = R$ gives $N = 10^{22}$. Putting all this together we arrive at the photon diffusion timescale:

$$t_{\text{diff}} = N t_{e\gamma} \quad (1.7)$$

which for the Sun is $\approx 30,000$ yr. This timescale is significantly shorter than the KH timescale, indicating that most of the Sun's thermal energy is stored in the random motions of electrons and ions rather than photons, even though the latter dominate the outward

transport of energy. Photons carry a larger fraction of the thermal energy of stars more massive than the Sun, as we will see in later lectures.

The relative magnitudes of these timescales are very different for most phases of evolution for most stars:

$$t_{\text{dyn}} \ll t_{\text{KH}} \ll t_{\text{nuc}} \quad (1.8)$$

This is important because it means that we can assume that stars are in hydrostatic and thermal equilibrium throughout most of their lives. The construction of accurate 1D stellar models is possible because of the disparity in these timescales. However, there are times when this assumption fails, e.g., during the end stages of massive stars.

Chapter 2

Radiation & Matter

Radiation (light) is the principle means by which we observe any phenomenon in the Universe (other observables include neutrinos, cosmic rays, and gravitational waves). Starlight provides us the opportunity to measure the detailed physical properties of stars. Furthermore, the interaction between radiation and matter plays an important role in the internal structure and evolution of stars. Much of this subject is covered in greater detail in AY200. This and the following lectures provide the necessary background for the rest of the course. The most important concepts in this lecture are the Stefan-Boltzmann Law, the Boltzmann distribution, and the Saha Equation.

We begin with a definition of the **specific intensity** (also known as the radiance):

$$I_\nu \equiv \frac{dE}{\cos \theta \, dA \, d\Omega \, d\nu \, dt} \quad (2.1)$$

which describes the amount of energy per unit time, per frequency, flowing through an area dA in a direction $d\Omega$, at an angle θ with respect to the unit normal. See Figure 2.1 for the relation between these geometric quantities. The specific intensity is the base quantity that we will use to derive other relevant aspects of the radiation field. In the absence of absorption or scattering, I_ν is *conserved*, which explains why we have defined a seemingly arbitrary and complicated quantity.

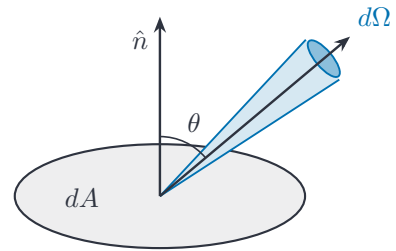


Figure 2.1: Geometry relevant for the definition of specific intensity.

Notation: You will often see descriptions of radiation expressed per unit frequency (e.g., I_ν) or per unit wavelength (e.g., I_λ). These two are related via: $I_\nu d\nu = I_\lambda d\lambda$. Relying on $c = \lambda\nu$, so $d\nu = \frac{c}{\lambda^2} d\lambda$, we find: $I_\nu = I_\lambda \frac{d\lambda}{d\nu} = \frac{\lambda^2}{c} I_\lambda$. Be careful of the units when you work with data, read papers, etc. There is no set convention, and quantities are often not well-labeled.

The **flux**, F , is defined as:

$$F_\nu = \frac{dE}{dA d\nu dt} \quad (2.2)$$

When I_ν is independent of θ , then $F_\nu = \pi I_\nu$. This is a case encountered frequently in astronomy. Notice that unlike I_ν , the flux depends on distance from the source ($\propto r^{-2}$). F_ν is often referred to as the monochromatic flux, as it is the flux at a specific frequency.

Radiation in thermodynamic equilibrium with matter has a distribution that depends only on temperature: $I_\nu = B_\nu(T)$, where T is the **kinetic temperature** of the matter. This “law” was proposed by Kirchhoff in 1860, and the associated radiation is referred to as **blackbody radiation**.

It was Planck who showed that $B_\nu(T)$ has the form:

$$B_\nu(T) = \frac{2h\nu^3/c^2}{e^{h\nu/kT} - 1} \quad (2.3)$$

and so the above equation is sometimes referred to as the Planck function, or Planck’s Law. In the above equation h is Planck’s constant and k is the Boltzmann constant. The key insight that lead to Planck’s Law was the realization that energy is quantized.

The **Rayleigh-Jeans limit** is found when setting ν to be small (or λ large). Noting that $e^x \approx 1 + x$ when $x \ll 1$, we find $B_\nu \approx \frac{2\nu^2 kT}{c^2}$, or $B_\lambda \propto \lambda^{-4}$. In this limit, the shape of the blackbody spectrum is independent of temperature — the temperature simply sets the normalization. This means that a measurement of the intensity provides a direct measurement of the temperature. In radio astronomy (where ν is small) this is referred to as the **brightness temperature**.

By taking the derivative of the Planck function we can derive a simple function relating the location of the peak of the function to temperature. This is known as **Wein’s Law**:

$$\lambda_{\max} T = 0.29 \text{ cm K} = 2900 \mu\text{m K} \quad (2.4)$$

This simple function is very valuable for estimating the wavelength at which an object radiating at temperature T will emit most of its light. Alternatively, if one measures the peak wavelength then one can estimate the temperature of the emitting source. For example,

the Sun has a temperature of 5800 K, which would imply $\lambda_{\max} \approx 0.5\mu m = 5000\text{\AA}$. It is of course not a coincidence that the peak is in the “visible” waveband. Use of Wein’s Law is a rather rough approximation for the temperature because, as we will see, real stellar flux distributions deviate from a blackbody function due to atomic and molecular absorption.

The **Stefan-Boltzmann Law** is derived by integrating F_ν over frequency:

$$F = \int F_\nu d\nu = \pi \int B_\nu d\nu = \sigma T^4 \quad (2.5)$$

where $\sigma = \frac{2\pi^5 k^4}{15h^3 c^2} \approx 5.67 \times 10^{-5} \text{ erg cm}^{-2} \text{ s}^{-1} \text{ K}^{-4}$ is the Stefan-Boltzmann constant. F without a subscript generally refers to the bolometric flux, i.e., the flux integrated over all frequencies (wavelengths). However, this is not always the case as subscripts are sometimes dropped for notational convenience (or carelessness). Be careful.

We can also define the flux as $F = L/A$ or luminosity per unit area. If we define $A = 4\pi R^2$ and insert the Stefan-Boltzmann Law for F we arrive at a very useful expression (please burn this one into your brain):

$$L = 4\pi R^2 \sigma T^4 \quad (2.6)$$

For completeness, we can define the radiation energy density as $u = aT^4$ which has units of erg cm^{-3} , where a is the radiation constant defined as $a \equiv 4\sigma/c$. The radiation pressure is $P_r = \frac{1}{3}u$. These quantities can be derived from the Planck function and basic thermodynamics.

It is very important to understand the state of matter as a function of temperature, density, and composition (number densities of all species). As we will see in later lectures, the opacity of matter is a sensitive function of its state. By “state” we mean the relative number densities of various species (different ionic states of a given element), whether atoms have combined to form molecules, as well as the distribution of excited states for each species and molecule. This is an enormously complicated task that is considerably simplified with a few key assumptions (in particular, LTE; see below). Two of the most important concepts are the Boltzmann distribution and the Saha Equation.

The **Boltzmann distribution**, $f(E) \propto e^{-E/kT}$, tells us the distribution of states of a system (e.g., of an ion), where each state has energy E .

Consider an atom with various level populations. The ratio of populations in two levels m and n is:

$$\frac{N_n}{N_m} = \frac{g_n}{g_m} e^{-\Delta E/kT} \quad (2.7)$$

where $\Delta E = E_n - E_m$ and g_i is the statistical weight (i.e., the number of states with the same energy). For example, for Hydrogen, $g_i = 2i^2$ where i is the shell number. An electron has $g_e = 2$ because there are two (degenerate) spin states at a given energy.

So, for a given temperature, the above equation tells us the distribution of *excitation states* for all atoms. This is very important! For example, if we can measure the relative populations of two states, we can infer the so-called **excitation temperature**. This is commonly done with observations of stellar spectra, in which the relative population of two levels from the same species (e.g., Fe II) can be determined.

We can express the populations of states relative to the *total* abundance of the ion in question by appealing to the partition function:

$$U(T) \equiv \sum_j g_j e^{-\Delta E_j/kT} \quad (2.8)$$

where the sum is over all j excited levels and ΔE_j is the energy difference relative to the ground state. The partition function is an intrinsic function of the ion (or molecule) in question and depends only on temperature. The abundance of ions in excited state j relative to the total number density of those ions is then:

$$\frac{n_j}{n} = \frac{g_j e^{-\Delta E_j/kT}}{U(T)} \quad (2.9)$$

The above arguments applied to the level populations of a given ion. By similar arguments (replacing level populations with ionization states), we can derive the **Saha Equation**. In equilibrium we have a balance between ionization and recombination: $\gamma + n_i \rightleftharpoons n_{i+1} + n_e$. As recombination generically scales as $n_{i+1}n_e$ (it is a two-body reaction), while ionization scales as n_i , we might expect: $\frac{n_i}{n_{i+1}n_e} = f(T)$ for some function f . You might guess that $f(T) \propto e^{-\Delta E/kT}$ where ΔE is the energy required to ionize. That would be a good guess!

The full derivation of the Saha Equation is tedious, but we can identify the fundamental dependencies with the following sketch. Consider an atom in its ground state with statistical weight g_0 . By some process the atom is ionized and produces an electron moving with momentum p . This would have required an energy $\chi + \frac{p^2}{2m_e}$ where χ is the ionization energy and m_e is the electron mass. The statistical weight of the final state is $g = g_{\text{ion}} g_e$ where g_{ion} is the statistical weight of the ion in the ground state and g_e is the number of possible states in which the free electron may be put. We know from quantum mechanics that every cell of phase space of volume h^3 may contain up to two electrons (one for each spin state). The number of cells available for the free electron with a momentum between p and $p + dp$ is $V_e 4\pi p^2 dp / h^3$ where $V_e = n_e^{-1}$ is the volume in ordinary space available per electron. We

can now express the ratio of ionized to neutral atoms via the Boltzmann equation (Eq. 2.7):

$$\frac{n_{i+1}}{n_i} = \frac{g_{i+1}}{g_i} \frac{2}{h^3 n_e} \int_0^\infty e^{-(\chi + p^2/2m)/kT} 4\pi p^2 dp \quad (2.10)$$

The integral can be solved analytically, yielding the Saha Equation:

$$\frac{n_{i+1} n_e}{n_i} = \left[\frac{2\pi m_e kT}{h^2} \right]^{3/2} \frac{2 g_{i+1}}{g_i} e^{-\Delta E / kT} \quad (2.11)$$

where we have re-written the ionization energy as ΔE . This equation defines the **ionization temperature**.

This is not the full story, as we have neglected excited states. The full treatment results in a replacement of the ground-state statistical weights with the partition functions of the n_{i+1} and n_i species.

Consider Hydrogen as an example. We will assume that all Hydrogen is in the ground state, so that $g_H = 2$. We also have $g_p = 1$. Furthermore, $n_{i+1} = n_p = n(\text{HII})$; $\Delta E = 13.6$ eV. This reduces to:

$$\frac{n(\text{HII})}{n(\text{HI})} \approx 2.4 \times 10^{15} \frac{T^{3/2}}{n_e} e^{-158,000/T} \quad (2.12)$$

where T is in K and n_e is in cm^{-3} .

Now lets compare to a typical “metal”, e.g., iron, as well as helium. The ionization energy of Fe I is 7.9 eV, while for He I it is 24.6 eV. This yields:

$$\frac{n(\text{FeII})}{n(\text{FeI})} \approx 2.4 \times 10^{15} \frac{T^{3/2}}{n_e} e^{-90,000/T}, \quad \frac{n(\text{HeII})}{n(\text{HeI})} \approx 2.4 \times 10^{15} \frac{T^{3/2}}{n_e} e^{-286,000/T} \quad (2.13)$$

Consider conditions at the Solar photosphere, where $T \approx 6000$ K. This leads to:

$$\frac{n(\text{FeII})}{n(\text{FeI})} \sim 10^5 \frac{n(\text{HII})}{n(\text{HI})}, \quad \frac{n(\text{HeII})}{n(\text{HeI})} \sim 10^{-10} \frac{n(\text{HII})}{n(\text{HI})} \quad (2.14)$$

This result is very important. In the Sun, the ratio of Fe to H atoms is $\approx 10^{-4.5}$. So even though Fe is *much* less abundant than H, because of the differences in ionization energies, there is more ionized Fe than ionized H in the solar photosphere. This implies that the Solar spectrum will be strongly affected by atomic Fe transitions, which in fact is what we observe.

Assuming a solar photosphere density of $n_e = 10^{13} \text{ cm}^{-3}$, we can compute the temperature at which the number density of ionized and neutral H atoms is equal. This occurs at 9000 K. For He, this occurs at 15,500 K. We can therefore expect to see few H I level transitions in the spectra of stars with $T \gtrsim 15,000$ K, and few He I transitions at $T \gtrsim 20,000$ K.

These results imply that there is a relatively narrow range of temperatures where electronic transitions of a given ionization state can occur. For example, at $T \lesssim 6000$ K all of the hydrogen is neutral and the energy of photons is too low to excite the $n = 2$ state, so we expect to observe no H lines. At higher temperatures all of the hydrogen will be ionized and so again no H lines will be observed.

These basic considerations go a long way toward explaining the observed spectral classification sequence.

Question: We can treat the formation and dissociation of molecules in a manner analogous to the Saha Equation for ions, replacing the ionization energy with the molecular dissociation energy. Most molecules have dissociation energies $\lesssim 5$ eV (CO is the most tightly bound molecule with a dissociation energy of 11 eV). What does this imply for the abundance of molecules at the surface of stars? In what types of stars do we expect molecules to be the most abundant?

Another very important concept is **Local Thermodynamic Equilibrium (LTE)**, which occurs when the following inter-related conditions are satisfied:

- The radiation field is equal to the Planck function, i.e., $I_\nu = B_\nu$.
- Level occupations follow the Boltzmann distribution for $T = T_{\text{kin}}$. This means that all of the temperatures defined above (kinetic, brightness, ionization, excitation) are the same.
- LTE occurs when collisions between the gas and photons are frequent enough to maintain equilibrium. This in turn implies that the opacity must be high.

LTE does *not* imply complete thermal equilibrium, otherwise there would be no radiative losses (e.g., at the stellar surface) and the temperature would be constant throughout the star. The relevant scale is the photon mean free path: if the temperature gradient is small over the mean free path length, then we can assume LTE.

Departures from LTE will occur when the gas density is low enough and/or the opacity is low enough so that departures from a Planck function occur. In this case the level populations will depend on the detailed radiation field and the problem becomes dramatically more complicated. For example, in the stellar interior we can assume LTE, so that the consideration of energy transfer requires simple integrals over the Planck function. In the outer

layers of the star where the densities are low, non-LTE (NLTE) becomes important (this is especially the case in metal-poor stars and the coronal region). In that case one must solve the detailed, frequency-by-frequency radiation transfer problem in order to understand the flow of radiation. NLTE effects are most important for strong lines in metal-poor stars. The computation of NLTE models is very challenging and is an active area of research. In this class we will deal exclusively with LTE conditions. If you are interested in NLTE, see the review by Asplund (2005).

Summary: Lets reflect on the importance of these results. A given stellar model will specify the temperature, density, and composition as a function of radius. Assuming LTE, we can use this temperature and the Saha Equation to determine the absolute densities of all atomic species and molecules. We can then use the Boltzmann distribution to specify the level populations of each species. The math becomes complicated when many elements are present as the equations become coupled through the electron density, but the physics is straightforward. As we will see in later lectures, the absolute abundance of each species determines the opacity of matter, both in the interior and at the photosphere. Opacity in the stellar interior affects the structure of the star, while opacity in the photosphere affects the emergent radiation (i.e., what we directly observe). This means that the solution to any stellar structure problem requires iteration: solve the equations, determine the state of matter and hence the amount of opacity, update the equations, repeat.

A key outstanding issue in this arena is the effect of NLTE. In some regimes this is the major limiting factor to obtaining more accurate physical properties of stars. Such models are becoming somewhat more widespread, but there are only 2-3 groups in the world working on this issue.

Chapter 3

Stellar Atmospheres

In this lecture we will cover the basics of radiative transfer, plane-parallel stellar atmospheres, and the formation and analysis of spectral lines. This topic could fill an entire course, so we cannot do the entire topic justice; our goal here is to understand the key concepts as they relate to observations of stellar surfaces.

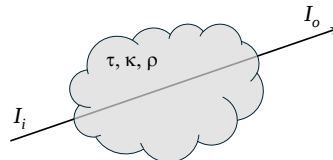
The **opacity** of a medium is an important concept that quantifies the probability of light interacting with matter. There are many different definitions of opacity (this is astronomy after all), so be aware of the units. Opacity is often written as κ with units of $\text{cm}^2 \text{g}^{-1}$. You can think of this as a cross section per unit mass.

The **optical depth** is an integral of the opacity and the density along a line of sight:

$$\tau \equiv \int \kappa \rho dz \quad (3.1)$$

along some dimension z . Notice that τ is unitless.

If we now consider an incident radiation field (I_i) passing through an absorbing medium, then the radiation field will be attenuated according to: $I_o = e^{-\tau} I_i$, or: $\frac{dI}{d\tau} = -I$ (see sketch).



In the general case, the opacity (and optical depth) will depend on wavelength; in such cases we write κ_ν and refer to the monochromatic opacity.

We may also encounter an emitting medium, quantified by j , the **emissivity** (also known as the emission coefficient). In this case we have $dI = j \rho dz$.

We can combine these two situations to arrive at a more general equation for monochromatic radiation transfer (in 1D):

$$\frac{dI_\lambda}{dz} = -\kappa_\lambda \rho I_\lambda + j_\lambda \rho \quad (3.2)$$

which quantifies the change in intensity, dI_λ , along a path length, dz , due to both absorption and emission.

This equation can be re-written as:

$$\frac{dI_\lambda}{d\tau_\lambda} = I_\lambda - S_\lambda \quad (3.3)$$

where we have redefined the optical depth so that $\tau = 0$ at the surface (e.g., of the star). S_λ is known as the **source function**. It is defined as $S_\lambda \equiv j_\lambda/\kappa_\lambda$, and represents the amount of emission per unit absorption. In LTE $S_\lambda = B_\lambda$.

The solution to this equation is obtained by use of an integrating factor:

$$\frac{d}{d\tau_\lambda}(I_\lambda e^{-\tau_\lambda}) = -S_\lambda e^{-\tau_\lambda} \quad (3.4)$$

Integrating:

$$I_\lambda e^{-\tau_\lambda} = I_{\lambda,0} e^{-\tau_{\lambda,0}} + \int_{\tau_\lambda}^{\tau_{\lambda,0}} S_\lambda e^{-\tau'} d\tau' \quad (3.5)$$

and rearranging:

$$I_\lambda = I_{\lambda,0} e^{\tau_\lambda - \tau_{\lambda,0}} + \int_{\tau_\lambda}^{\tau_{\lambda,0}} S_\lambda e^{\tau_\lambda - \tau'} d\tau' \quad (3.6)$$

this allows us to calculate the radiation field once we know the source function, S_λ as a function of depth. This looks relatively straightforward, but in fact in many situations S_λ depends on opacity and I_λ .

The stellar atmosphere represents the outer layers of a star and comprises the transition between the optically-thick ($\tau \gg 1$) and optically-thin ($\tau < 1$) regions. It is in the atmosphere that the stellar spectrum is formed (specifically in the photosphere). The very outer layers of the atmosphere contain the chromosphere and corona.

The atmospheres of most stars are very thin compared to the stellar radius, and so it is usually a good assumption to treat the atmospheric structure as plane-parallel, or 1D, in which the physical variables depend only on vertical height, z . This assumption breaks down for giant and supergiant stars, which have very extended atmospheres. We're now going to work through some of the basic features of **plane-parallel atmospheres**. The extension to spherical geometry is straightforward but more mathematically cumbersome.

We begin by defining $\tau_\lambda = \tau_{\lambda,v} / \cos \theta$ where $\tau_{\lambda,v}$ is the optical depth along a path normal to the surface (i.e., vertical), and τ_λ is the optical depth at an angle θ with respect to the normal (see the figure at right for the geometry).

Our radiative transfer equation becomes: $\alpha \frac{dI_\lambda}{d\tau_{\lambda,v}} = I_\lambda - S_\lambda$ with $\alpha \equiv \cos \theta$. At the surface we have $I_\lambda(\tau_{\lambda,v} = 0) = 0$ for $\alpha < 0$. In other words, at the surface there is no radiation coming down from above. This insight when applied to Eq. 3.6 allows us to compute the emergent specific intensity:

$$I_\lambda(\tau_{\lambda,v} = 0, \alpha > 0) = \int_0^\infty S_\lambda e^{-\tau'/\alpha} \frac{d\tau'}{\alpha} \quad (3.7)$$

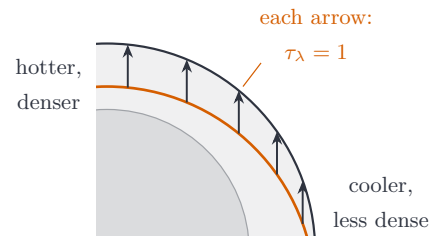
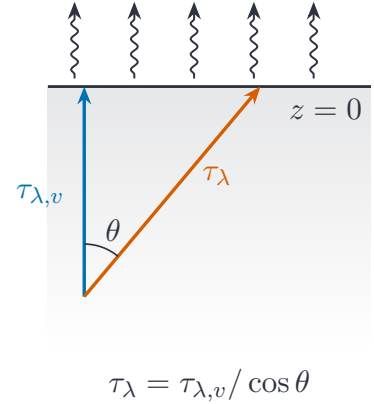
Specifically, this equation represents the intensity emitted at the surface ($\tau = 0$), which is determined by a weighted sum of the source function over all depths.

Near the surface we can expand S_λ as a Taylor series: $S_\lambda(\tau_{\lambda,v}) = S_\lambda(0) + S'_\lambda(0) \tau_{\lambda,v} = a + b \tau_{\lambda,v}$. If we insert this equation into Eq. 3.7 we find: $I_\lambda(\tau_{\lambda,v} = 0, \alpha > 0) = a + b\alpha$, but this is just equal to the source function at $\tau_{\lambda,v} = \alpha$. i.e., $I_\lambda(\tau_{\lambda,v} = 0, \alpha > 0) = S_\lambda(\tau_{\lambda,v} = \alpha)$. Recalling our definition of $\tau_\lambda = \tau_{\lambda,v} / \alpha$ we arrive at an important equation:

$$I_\lambda(\tau_{\lambda,v} = 0, \alpha > 0) = S_\lambda(\tau_\lambda = 1) \quad (3.8)$$

In words: *the emergent specific intensity at the surface is equal to the source function at optical depth equal unity*. This should make sense - the photons we see from the Sun were, on average, emitted at a depth in the atmosphere corresponding to $\tau = 1$. If photons were emitted any deeper then there would be a high probability of being absorbed/scattered again. The equation above is known as the **Eddington-Barbier relation**. There are several important implications of this relation that we now discuss.

Limb Darkening. In a star, the source function generally increases with increasing τ (because the temperature increases with increasing depth). This means that a vertical ray ($\alpha = 1$) will emerge at $\tau_\lambda = \tau_{\lambda,v} = 1$. However, a (nearly) horizontal ray, e.g., at $\alpha = 0.1$, will emerge at $\tau_\lambda = \tau_{\lambda,v} / 0.1 = 1$, i.e., $\tau_{\lambda,v} = 0.1$, which is much higher in the atmosphere (see sketch). Higher in the atmosphere the temperature is lower, so the source function (and hence the flux) is lower. This means that when you look at the “edge” (limb) of the Sun, it will appear darker than the center. This is called



limb darkening, and it is an important probe of the temperature profile in the photosphere (see the sketch at right; each arrow corresponds to $\tau_\lambda = 1$). In addition to limb darkening, the limb will also appear redder, as the outer layers are cooler.

A second (and much more important) application of the Eddington-Barbier relation is an understanding of the formation of **spectral lines**. An absorption line is formed when κ_λ is larger than average over a narrow wavelength range. In this line $\tau_\lambda = 1$ occurs much higher in the atmosphere, where T is lower and hence $B_\lambda(T)$ is lower in the line. Lower flux in the line corresponds to a darker appearance, i.e., an absorption line. Notice that the appearance of dark absorption lines is *not* simply a consequence of more absorption at that wavelength - if the relation between temperature and optical depth was flat then the line would not be dark. *It is the depth-dependence of the source function that gives rise to spectral lines.* As such, detailed analysis of spectral lines offers insight into the temperature profile in stars.

Question: You observe a spectrum of an M-type dwarf star and find, to your surprise, that the H α line is *in emission*. Suppose I told you that the line-forming region for H α (where $\tau_{\text{H}\alpha} = 1$) is in the chromosphere. What could you infer about the properties of the chromosphere from these facts?

Models for stellar atmospheres form the backbone of many areas of research, including interpreting observations of brown dwarfs and exoplanets, and deriving detailed properties of stars based on observed spectra. They are even essential for the interpretation of spectra of distant galaxies. One of the most popular programs for computing stellar atmospheres is **ATLAS** and was created by Bob Kurucz here at the CfA beginning in the 1960s (other popular programs include **MARCS**, **PHOENIX**, **TLUSTY**, and **CMFGEN**).

At its core, a model atmosphere specifies the temperature and pressure profiles in a star, e.g., $T(\tau)$ and $P(\tau)$. How would we compute these quantities? First, we must assume an equation of state (EOS), which sets the relation between T , P , and ρ (e.g., can assume an ideal gas). Second, it is common to assume hydrostatic equilibrium (more on that below), and the sources of opacity, which will in general depend on T and P , so $\kappa_\lambda(T, P)$. Finally, we employ the radiative transfer equation to determine the flow of radiation through the atmosphere. The full problem is extremely complex because everything depends on everything else. For example, if you change the temperature, then the pressure will change, as will the opacity, which means the radiation flow will change, which in turn affects the temperature. So the problem must be solved in an iterative fashion.

How far can we get with simple models? In HW1 you will derive simple relations for $T(\tau)$ and $P(\tau)$ that are surprisingly accurate and which illuminate the basic physics involved.

Hydrostatic Equilibrium (HE) is a key concept in many areas of astrophysics in which a gas is supported against gravity by pressure gradients. In 1D we have:

$$\frac{dP}{dz} = -g\rho, \quad g = \frac{GM}{R^2} \quad (3.9)$$

where g is the surface gravity. Since the atmosphere is thin compared to the stellar radius, we can assume that g is constant. So g , or $\log g$, is a useful way of describing an atmosphere. It is for this reason that stellar atmospheres are characterized by a temperature and $\log g$.

For an ideal gas we have $\rho = P\mu/kT$, where μ is the mean molecular weight (discussed in detail in Lecture 4). We can therefore re-write the HE equation as:

$$\frac{dP}{dz} = -g \frac{P\mu}{kT} \quad (3.10)$$

For an isothermal atmosphere ($T = \text{constant}$) the solution is: $P = P_0 e^{-z/H_p}$ where $H_p = \frac{kT}{g\mu}$ is the pressure scale height. For the Sun $H_p = 300 \text{ km}$, or $0.0004R_\odot$ – the atmosphere is very thin indeed. We can write the HE equation one more way by recalling that $d\tau = -\kappa \rho dz$ (Eq. 1), which leads to $dP/d\tau = g/\kappa$.

A **grey atmosphere** is one in which the opacity is independent of wavelength. This assumption allows for analytic computation of model atmospheres, and gets the basic behavior approximately correct.

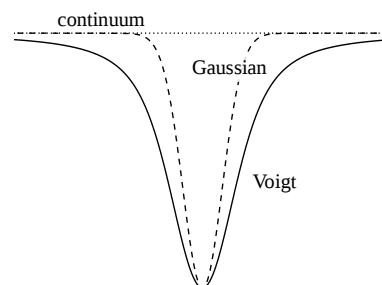
The **Rosseland mean opacity** is defined as:

$$\kappa_R = \left[\frac{\int d\nu \kappa_\nu^{-1} \partial B / \partial T}{\int d\nu \partial B / \partial T} \right]^{-1} \quad (3.11)$$

this is a harmonic mean in which the greatest contribution comes from the lowest values of the opacity. See Lecture 7 (or Ch 5.1 in Lamers & Levesque) for discussion of why one would define a mean opacity in this way. Note that τ_R is frequently used as the depth variable in stellar atmosphere models.

Spectral lines are due to electronic transitions within atoms, or electronic, vibrational, or rotational transitions within molecules. In general, atomic transitions are stronger than molecular transitions, but there are many more degrees of freedom in molecules, so there are many more total transitions available in molecules.

The **equivalent width** (EW, or W) is the integral of the line depth relative to the continuum. The EW often has



units of wavelength. Note that line strengths are always measured relative to a continuum (see sketch), so the EW depends on both the line strength and the source of continuum photons. This sounds like a technical detail, but it is important!

The detailed shape of a spectral line is known as the **line profile** and is affected by multiple physical processes:

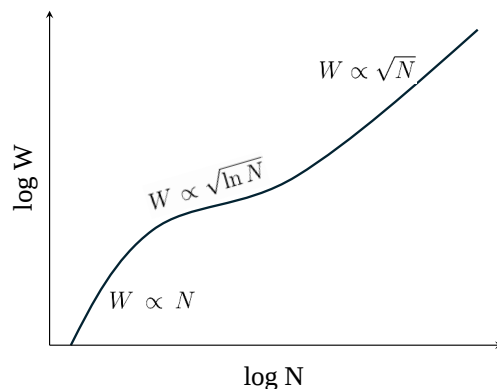
- natural line width: this is the minimum line width allowed by quantum mechanics and it is set by the uncertainty principle in the form: $\Delta E \Delta t = \hbar$ where Δt refers to the lifetime of an excited state. The line profile associated with this process can be approximated by the Lorentzian profile, which is $f(\lambda) \propto \frac{\gamma}{a(\lambda - \lambda_0)^2 + \gamma^2}$ where γ is the line strength, λ_0 is the line center, and a is a constant.
- pressure broadening: a result either of direct collisions or electric field effects (Stark broadening). In the former case, direct collisions shorten the characteristic lifetime of the excited state, which via the uncertainty principle results in a larger uncertainty in energy, hence larger line width. The Stark effect, which is important for Hydrogen lines, describes the change in energy levels due to the presence of an electric field.
- doppler broadening: due to thermal motions of the particles. Atoms move with a speed, v , that is related to their thermal energy by: $\frac{1}{2}mv^2 = kT$. For H at $T = 5,000\text{K}$ this corresponds to a velocity of 9 km s^{-1} while for Fe at the same temperature the velocity is 1 km s^{-1} .
- microturbulence: non-thermal, small-scale motion within the star (e.g., driven by turbulence). “Small” means that the scale of the turbulence is smaller than the mean free path of the photon. Microturbulence will both broaden and strengthen the line (e.g., the EW increases). Line strengthening due to microturbulence is strongest for lines that are saturated (i.e., on part II of the curve-of-growth; see below).
- macroturbulence: due to large-scale motions within the star (e.g., convective cells). “large” means that the scale of motion is larger than the mean free path of the photon. Macroturbulence will broaden a line, but not change its EW.
- rotational broadening: due to stellar rotation.

The first three broadening terms above are usually self-consistently predicted in models of stellar atmospheres. The latter three are not. In principle one can separately measure them

in data, but in practice that requires extremely high signal-to-noise ratios (SNR) and instrument spectral resolution. 3D simulations of stellar atmospheres that resolve the convective cells do make self-consistent predictions for the micro and macroturbulence. Stellar rotation is an “extrinsic” variable that cannot be predicted from first principles.

The combined spectral line profile is well described in most cases by the Voigt profile, which has a Gaussian core and Lorentzian wings (see sketch). The Gaussian core is governed by the thermal motions of the absorbers, while the Lorentzian profile reflects the physics governing the natural line width as well as pressure broadening.

Starting from basic considerations about the relation between line depths and the column density of absorbers, N , one can derive relations between N and the line EW (W). Such a relation is known as the **curve of growth** (see sketch), and has three regimes. In regime I the lines are weak, the profiles are mostly Gaussian, and $W \propto N$. When the column density of absorbers is large enough the line becomes saturated. This is regime II, where $W \propto \sqrt{\ln N}$. Saturated lines are not useful for determining abundances of species (a large change in species abundance produce a small change in observed EW). In regime III the damping wings become prominent and the EW again increases with N , though more slowly than before: $W \propto \sqrt{N}$. Supplementary material at the end of this lecture contains additional details regarding the curve of growth.



Spectral resolution refers to how finely one can resolve individual spectral features. One can quantify the spectral resolution by the dimensionless parameter R , defined as $R \equiv \lambda/\Delta\lambda$ where $\Delta\lambda$ is the full-width at half maximum (FWHM). Be aware that $\text{FWHM} \approx 2.35\sigma$ where σ is the standard deviation of a Gaussian. Sometimes the resolution is expressed in terms of σ , or $\Delta\lambda$, or by expressing the above in terms of velocity via: $\Delta v = c/R$ where c is the speed of light. Resolution is not the same as *sampling*. The latter determines how finely one samples (e.g., in wavelength space) a fixed line profile. Nyquist sampling is when the sampling is $2\times$ the resolution. Anything less than Nyquist is referred to as “under-sampling” and should be avoided, as information is lost in an under-sampled spectrum.

Astronomical objects have an “intrinsic” resolution set by the characteristic velocity of the system. The Sun has surface motions of order 1 km s^{-1} , which sets a maximum useful resolution at $R \approx 300,000$. In contrast, galaxies have internal motions of order 100 km s^{-1} , setting a maximum useful spectrograph resolution at $R \approx 3,000$.

Depending on the science case of interest, one might design a high ($R \sim 10^5$), moderate ($R \sim 10^4$) or low ($R \sim 10^3$) resolution spectrograph. The cost of a spectrograph will scale as (among other things) the number of pixels required at the detector. For a fixed number of pixels, higher resolution means a smaller wavelength range. The design of spectrographs must therefore carefully consider trade-offs between resolution, wavelength range, and multiplexing (number of sources observed at one time).

Summary: When we observe a spectrum of any star, we are observing the $\tau \approx 1$ surface at each wavelength. The combination of opacity effects and a depth-dependent source function produces a spectrum full of complex absorption lines. The spectral lines are sensitive to the physical properties of the star, including its effective temperature, T_{eff} , surface gravity, $\log g$, and composition. The line shapes are sensitive to additional effects including surface rotation, convective motions, and the presence and strength of magnetic fields (via Zeeman splitting). It is a major industry to use radiative transfer models to infer the physical properties of stars based on their observed spectra.

There are many areas of active research in this field: 1D plane-parallel atmospheres are giving way to realistic 3D models with self-consistent convection and magnetic fields. As mentioned in the previous lecture, relaxing the assumption of LTE adds considerable complexity but appears to be necessary in at least some situations (particularly for metal-poor stars). The atomic and molecular parameters necessary to construct realistic stellar spectra are in many cases poorly known. Progress can be made here by increasingly accurate quantum mechanical models of atoms (e.g., the Hartree-Fock method) and laboratory measurements. These approaches suffer from poor funding and a dearth of astronomers trained in the relevant techniques.

Supplementary Material

The following provides a fairly detailed exploration of the three regimes in the curve of growth (following Draine 2011, Ch 9).

Let’s consider a single absorption line with central wavelength λ_0 (and corresponding central

frequency ν_0). If the line is relatively well-isolated from other lines then we can identify the local stellar continuum flux adjacent to the line and use that information to define, via interpolation, a continuum flux even at wavelengths where the line is present (see sketch associated with our first mention of equivalent width earlier in this lecture). Define this continuum flux as $F_\nu(0)$. In most applications relevant for stellar atmospheres we can make the assumption that:

$$F_\nu = F_\nu(0)e^{-\tau_\nu} \quad (3.12)$$

where τ_ν is the frequency (wavelength)-dependent optical depth.

For a transition between a lower (l) and upper (u) electronic state, the optical depth is related to the column density of absorbers in the lower state, N_l , via:

$$\tau_\nu = \frac{\pi e^2}{m_e c} f_{lu} N_l \phi_\nu = \tau_0 \phi_\nu \quad (3.13)$$

where f_{lu} is the **oscillator strength** and is related to the Einstein A coefficient. The oscillator strength is an intrinsic property of the transition and is either provided by laboratory measurements or ab initio quantum mechanical calculations. The term ϕ_ν is the normalized line profile and is usually represented as a Voigt profile, which has a Gaussian core and Lorentzian wings.

We can then define the dimensionless equivalent width:

$$W \equiv \int \frac{d\nu}{\nu_0} \left[1 - \frac{F_\nu}{F_\nu(0)} \right] = \int \frac{d\nu}{\nu_0} (1 - e^{-\tau_\nu}) \quad (3.14)$$

where the “0” subscript refers to quantities at line center.

We will now consider three regimes: optically-thin, saturated, and damped absorption.

Optically-thin absorption In this regime absorption is small, so $\tau_0 \lesssim 1$. If we assume that the velocity distribution of the absorbers is Gaussian with 1D dispersion σ_V , then the optical depth near line center can be written as:

$$\tau = \tau_0 e^{-(v/b)^2} \quad (3.15)$$

where the velocity, v is relative to line-center ($dv/c = d\nu/\nu_0$) and $b \equiv \sqrt{2}\sigma_V$. The parameter b is known as the **microturbulence**. The optical depth at line-center is then

$$\tau_0 = \sqrt{\pi} \frac{\pi e^2}{m_e c} \frac{f_{lu} N_l \lambda_0}{b} \quad (3.16)$$

Since we are assuming τ is small, we may approximate $(1 - e^{-\tau}) \approx \tau$ in Eqn. 3.14, and obtain:

$$W \approx \sqrt{\pi} \frac{b}{c} \tau_0 \quad (3.17)$$

which leads directly to the result that when lines are weak the equivalent width scales linearly with the column density of the absorber: $W \propto N_l$. This is regime I in the curve of growth.

Saturated absorption In this regime the core of the line depth is saturated and so no longer increases as the number density of absorbers increases. In this case the absorption line can be approximated as box-shaped with the width governed only by the doppler broadening of the line. In this case the equivalent width can be approximated as:

$$W \approx \frac{(\Delta\nu)_{\text{FWHM}}}{\nu_0} = \frac{(\Delta v)_{\text{FWHM}}}{c} \approx \frac{2b}{c} \sqrt{\ln(\tau_0/\ln 2)} \quad (3.18)$$

where FWHM is the full width at half maximum. Notice that in this case W depends linearly on b , which means the equivalent width is sensitive to the small-scale motions of atoms in the photosphere. In contrast, W is very weakly dependent on N_l : $W \propto \sqrt{\ln N_l}$. It is for this reason that this regime is referred to as the flat portion of the curve of growth. Abundances of elements therefore cannot be inferred from saturated lines.

Damped absorption In this regime the Doppler core is completely saturated but the damping wings of the line become sufficiently strong to affect the total equivalent width. The damping wings are the result of natural line broadening and are therefore well-described by the Lorentzian profile:

$$\tau_\nu = \frac{\pi e^2}{m_e c} f_{lu} N_l \frac{4\gamma_{lu}}{16\pi^2(\nu - \nu_0)^2 + \gamma_{lu}^2} \quad (3.19)$$

where γ_{lu} is related to the Einstein A coefficients and is, like the oscillator strength, a fundamental property of the atomic transition.

Approximating the equivalent width via the FWHM, we have:

$$W \approx \frac{\nu - \nu_0}{\nu_0} \quad (3.20)$$

Because we are considering the damping wings, $(\nu - \nu_0)^2$ is much larger than γ_{ul} , simplifying the denominator in the Lorentzian. After some algebra we arrive at the conclusion that the equivalent width in this regime scales as $W \propto \sqrt{N_l}$. A challenge of measuring element abundances from saturated lines is that the damping wings span a large range in wavelength and will therefore often be affected by absorption from other lines. In order to extract information from damped lines one must therefore simultaneously model a suite of atomic transitions from different elements.

Chapter 4

Equation of State

We are now going to turn our attention to the physical principles governing stellar structure and evolution. In this lecture and the next we will review the essential “microphysics” necessary to construct stellar models. These include the equation of state, sources of opacity, and nuclear energy generation.

The **equation of state** (EOS) is the general relation between pressure, density, and temperature, e.g., $P = P(\rho, T)$. The EOS is in detail very complicated and depends on the composition of the gas, whether the particles are relativistic, and whether the radiation comprises a significant fraction of the total pressure. In this lecture we will cover the most common limiting cases, as it turns out that regions of stars are frequently very well-characterized as being in one particular EoS limit. From a practical point of view, the EOS is necessary to “close” the equations of stellar structure, as we will see in Lecture 8. Much of the material in this lecture should be familiar from your undergraduate physics courses.

Recall from quantum mechanics that an elementary cell in phase space has the volume: $\Delta x \Delta y \Delta z \Delta p_x \Delta p_y \Delta p_z = h^3$ where $h = 6.63 \times 10^{-27}$ erg s is Planck’s constant. The number of particles per phase space cell is

$$n = \frac{g}{e^{(E-\mu)/kT} \pm 1} \quad (4.1)$$

where g is the number of states within the cell, E is the particle energy, μ is the chemical potential¹, and $+$ for fermions (non-integer spin, e.g., electrons or protons), and $-$ for bosons

¹If you find the concept of the chemical potential opaque, you are not alone. In classical thermodynamics one can define the chemical potential in several ways, including: $\mu \equiv (\partial E / \partial N)_{S,V}$, i.e., the change in energy due to changing the number of particles at constant entropy and volume. For an extended discussion of the chemical potential in various contexts, see Baierlein (2001).

(integer spin, e.g., photons). When the above equation has $+$ or $-$ it is known as a **Fermi-Dirac** or **Bose-Einstein** distribution.

The total energy, E , can be written as $E = mc^2 + E_k$ where E_k is the kinetic energy. Appealing to special relativity we can write the total energy as $E^2 = (mc^2)^2 + (pc)^2$ where p is the momentum. In the non-relativistic (NR) limit ($p \ll mc$) $E_k = p^2/2m$, while in the ultra-relativistic (UR) limit ($p \gg mc$) $E_k = pc$. This implies that $v = p/m$ (NR) and $v = c$ (UR).

The number density of particles (per cm^3) with momenta between p and $p + dp$ is given by

$$n(p)dp = \frac{g}{e^{(E-\mu)/kT} \pm 1} \frac{4\pi p^2}{h^3} dp \quad (4.2)$$

The volume term arises by considering a spherical shell in momentum space with surface $4\pi p^2$ and thickness dp . The number density per volume integrating over all momenta is then simply $n = \int n(p) dp$. One can convert to mass density via $\rho = nm$.

The kinetic energy of all particles in a unit volume can be calculated as $K = \int E_k(p) n(p) dp$, which has units erg cm^{-3} .

Pressure (P) can be defined as the momentum flux across a surface, integrated over all particles. For an isotropic gas we can select the unit surface to be perpendicular to the x -axis, and we may then calculate the pressure as:

$$P = \int v_x p_x n(p) dp = \frac{1}{3} \int v(p) p n(p) dp = \frac{1}{3} n \langle vp \rangle \quad (4.3)$$

where the second expression arises from considering: $\mathbf{v} \cdot \mathbf{p} = vp = v_x p_x + v_y p_y + v_z p_z$ and from isotropy we have $\langle v_x p_x \rangle = \langle v_y p_y \rangle = \langle v_z p_z \rangle$ so $\langle v_x p_x \rangle = \frac{1}{3} \langle vp \rangle$.

Notice that in the NR-limit we have $K = \frac{3}{2} P$ and in the UR-limit we have $K = 3 P$. These expressions are very general and do not depend on the distribution function $n(p)$. In deriving these equations we have made the implicit assumption that the particles are interacting so weakly that the energy of interactions may be neglected, but strongly enough to establish an equilibrium distribution.

We are now going to embark on an extended digression on the relation between ρ and n , which depends on the composition of the material.

Recall that we define X , Y , and Z as the fractional abundance (by mass) of H, He, and all heavier elements, and that $X + Y + Z = 1$. The Sun has approximate values of $X_\odot = 0.7$, $Y_\odot = 0.28$, and $Z_\odot = 0.02$ (there is considerable controversy regarding the solar composition

– some studies advocate $Z_\odot = 0.014$). Our goal here is to form expressions relating ρ and n that depend on X , Y , and Z . In the following we will assume a gas that is fully ionized, which is a good assumption in most stellar applications.

Recall that the mass of one H atom is $m_H = 1.67 \times 10^{-24}$ g. The mass of a helium atom is (not exactly but close to) $m_{He} = 4m_H$, and the mass of element Z_p is approximately $2Z_p m_H$. The average proton number of heavy elements is $Z_p = 8$, because it so happens that oxygen is the most abundant heavy element. The mass of heavy elements can therefore be reasonably approximated by $16m_H$.

With these approximations we can estimate the number density of ions as

$$n_i = \frac{\rho}{m_H} \left(X + \frac{Y}{4} + \frac{Z}{16} \right) \quad (4.4)$$

For a fully-ionized gas we have 1 e^- per nucleon for H, 1 e^- per 2 nucleons for He, and approximately 1 e^- per 2 nucleons for heavier elements. So the number density of electrons is just:

$$n_e = \frac{\rho}{m_H} \left(X + \frac{Y}{2} + \frac{Z}{2} \right) = \frac{\rho}{m_H} \frac{1+X}{2} \quad (4.5)$$

(remember that Z is the mass fraction of heavy elements, while Z_p is the nuclear charge). The total number density of particles is therefore

$$n = n_i + n_e = \frac{\rho}{m_H} (2X + 0.75Y + 0.5Z) \quad (4.6)$$

Now we can define the **mean molecular weight** as the average mass of a particle in units of the mass of hydrogen:

$$\mu \equiv \frac{\rho}{n m_H} \approx \frac{1}{2X + 0.75Y + 0.5Z} \quad (4.7)$$

$$\mu_i \equiv \frac{\rho}{n_i m_H} = \frac{1}{X + Y/4 + Z/16} \quad (4.8)$$

$$\mu_e \equiv \frac{\rho}{n_e m_H} = \frac{2}{1 + X} \quad (4.9)$$

note that the mean molecular weight has nothing to do with the chemical potential, even though they share the same variable (yes, it is confusing).

As an example, these definitions allow us to write the ideal gas pressure as:

$$P_i = \frac{\rho k T}{\mu_i m_H} \quad , \quad P_e = \frac{\rho k T}{\mu_e m_H} \quad , \quad P = P_e + P_i = \frac{\rho k T}{m_H} \left(\frac{1}{\mu_e} + \frac{1}{\mu_i} \right) \quad (4.10)$$

where $\mu^{-1} = \mu_e^{-1} + \mu_i^{-1}$

These relations highlight a very important point: for a given total mass or density profile in a star, the composition will affect the pressure and therefore both the stellar structure and evolution. This partly explains why composition is a fundamental variable in the construction of stellar models.

Question: How else might composition affect the structure and evolution of a star?

We will now consider several important limiting cases for the EOS. The basic options are: degenerate or not, relativistic or not, radiation or not.

Radiation Pressure: Let's consider the distribution function for photons (bosons). In this case we have $E = h\nu$, $p = h\nu/c$, $\mu = 0$ (chemical potential), and $g = 2$ (two polarization states). With these ingredients we can write the Bose-Einstein distribution as:

$$n_\nu d\nu = \frac{8\pi}{c^3} \frac{\nu^2 d\nu}{e^{h\nu/kT} - 1} \quad (4.11)$$

which is just the Planck function. Integrating this distribution function to compute the pressure and energy densities leads to an energy density in radiation of $u_r = aT^4$ where a is the radiation constant, and radiation pressure of $P_r = \frac{1}{3}u_r$.

In general, $P = P_{\text{gas}} + P_r$. For solar composition these two pressure terms are equal when $\rho = \left(\frac{T}{3 \times 10^7 \text{ K}}\right)^3 \text{ g cm}^{-3}$. So at very high temperatures, this threshold occurs at densities relevant for stellar interiors. The importance of radiation pressure for mechanical stability will become clear in later lectures.

Ideal Gas Pressure: An ideal gas can be thought of as a non-relativistic classical (non-degenerate) gas. These conditions are satisfied when $mc^2 \gg E - \mu \gg kT$. In this limit the “+1” factor in the distribution function can be ignored as the exponential term in the denominator of the distribution function is very large. This means that $n \ll 1$ and there are few particles per elementary phase space cell. In this limit we have:

$$n(p) \propto p^2 e^{-p^2/2mkT} \propto v^2 e^{-mv^2/2kT} \quad (4.12)$$

which is just the **Maxwellian distribution**. If we plug this distribution function into our equations for n and P we arrive at the classic EOS for an ideal gas: $P = nkT$.

Degeneracy Pressure: We will now consider the EOS for degenerate fermions. Fermions

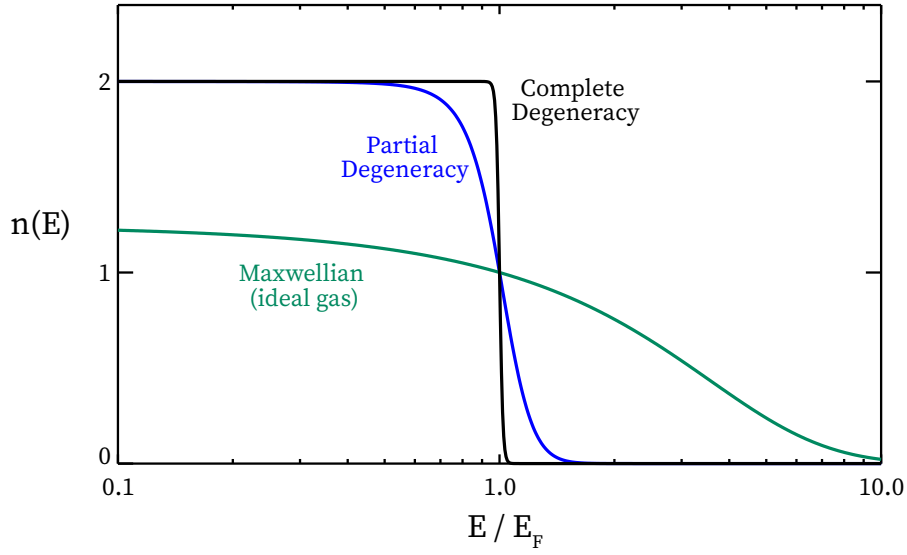


Figure 4.1: Number of particles per phase space cell as a function of energy (E), plotted in units of the Fermi Energy (E_F). The ideal gas, partial degeneracy, and complete degeneracy solutions are plotted.

are governed by the Fermi-Dirac distribution function:

$$n = \frac{2}{e^{(E-\mu)/kT} + 1} = \frac{2}{e^{(E_k-E_F)/kT} + 1} \quad (4.13)$$

where E_F is the Fermi energy (and p_F is the Fermi momentum). The second expression arises because the chemical potential $\mu = mc^2 + E_F$. Loosely speaking, the Fermi energy is the highest energy occupied in a degenerate gas. One can see this by considering the limit $E_F \gg kT$. In this case $n = 2$ for $E_k < E_F$ and $n = 0$ for $E_k > E_F$. When the temperature drops, particles occupy the lowest energy states available. However, fermions obey Pauli's exclusion principle: one particle per quantum state. Particles therefore fill states up to the Fermi energy. A gas that behaves this way is called degenerate. See Fig. 4.1 for examples.

We can see that the Fermi energy (momentum) is related to the density of particles by integrating the Fermi-Dirac distribution (assuming $E_F \gg kT$):

$$n = \int n(p) dp = 2 \int_0^{p_F} \frac{4\pi p^2}{h^3} dp = \frac{8\pi}{3} \left(\frac{p_F}{h} \right)^3 \quad (4.14)$$

If we compute the pressure using the F-D distribution with an upper-limit on the momentum integration of p_F (expressed in terms of n via Eq. 4.14), we arrive at two limiting cases:

$$P \propto \rho^{5/3}, \quad p \ll mc \quad (\text{non-relativistic}) \quad (4.15)$$

$$P \propto \rho^{4/3}, \quad p \gg mc \quad (\text{ultra} - \text{relativistic}) \quad (4.16)$$

In the NR limit $E = p^2/2m$ so $E_F = \frac{\hbar^2}{2m} \left(\frac{3n}{8\pi}\right)^{2/3}$. Note that $m_p/m_e \approx 1836$ (ratio of proton to electron mass), so $E_F(e^-) \gg E_F(p)$ at a given number density. This means that electrons will become degenerate before ions. We will encounter situations in which pressure support is provided predominantly by degenerate electrons (e.g., white dwarfs, degenerate helium cores of low-mass stars) and degenerate neutrons (neutron stars).

The general expression for the pressure of a cold degenerate electron gas is:

$$P_e = \frac{8\pi}{3} \frac{m_e^4 c^5}{h^3} \int_0^{x_F} \frac{x^4 dx}{(1+x^2)^{1/2}} = A f(x) \quad (4.17)$$

with $A = 6.002 \times 10^{22}$ dyne cm^{-2} for electrons and

$$f(x) = x(2x^2 - 3)(1+x^2)^{1/2} + 3 \sinh^{-1} x \quad (4.18)$$

where $x = \left(\frac{\rho}{B\mu_e}\right)^{1/3}$, $B = 9.739 \times 10^5$ g cm^{-3} and $\mu_e = 2/(1+X)$.

You will often hear talk of “stiff” and “soft” equations of state. A stiff EOS is one in which $d \ln P / d \ln \rho$ is large, so that a small change in density results in a large change in pressure (e.g., metal has a stiff EOS, whereas a sponge has a soft EOS).

Summary: There are four main sources of pressure support that will be relevant for stars: ideal gas ($P_{\text{gas}} = nkT$), radiation ($P_r = \frac{1}{3}aT^4$), degenerate, non-relativistic gas ($P_{\text{deg,NR}} \propto \rho^{5/3}$), and degenerate, ultra-relativistic gas ($P_{\text{deg,UR}} \propto \rho^{4/3}$). In nearly all regimes of interest the pressure is dominated by one of these relations. Summaries of the various regimes in which certain EOS's dominate are shown in Figures 1 and 2.

The state of the art in this field involves complex calculations tailored to a specific subset of the temperature-density plane. For example, the EOS for low-temperature matter requires good knowledge of molecule formation, while an EOS at high temperatures and/or densities requires inclusion of pair creation and various non-ideal effects. At the very highest densities relevant for compact remnants (e.g., neutron stars) the EOS is highly uncertain. The reason for this is that ab initio QCD calculations do not exist for the thermodynamic conditions relevant for neutron stars. This means that astronomical observations of these objects can provide novel constraints on the EOS of dense matter.

Supplementary Material

We often encounter situations where a process occurs on a timescale much shorter than the time it takes for heat exchange to occur. Such a process is called **adiabatic**. In this section we provide a brief overview of thermodynamic quantities and adiabatic processes.

A combination of the first and second laws of thermodynamics leads to the following important expression:

$$\delta q = Tds = du + Pdv = du - \frac{P}{\rho^2}d\rho \quad (4.19)$$

where δq is the change in heat content, du is the change in internal energy (per unit mass), s is the specific entropy (entropy per unit mass) and $v \equiv 1/\rho$ is the volume of a unit mass.

We can write the EOS in a general form as:

$$\frac{dP}{P} = \chi_T \frac{dT}{T} + \chi_\rho \frac{d\rho}{\rho} \quad (4.20)$$

where

$$\chi_T \equiv \left(\frac{\partial \ln P}{\partial \ln T} \right)_\rho \quad (4.21)$$

and

$$\chi_\rho \equiv \left(\frac{\partial \ln P}{\partial \ln \rho} \right)_T \quad (4.22)$$

For an ideal gas we have $\chi_T = \chi_\rho = 1$ while for radiation-dominated gas we have $\chi_T = 4$ and $\chi_\rho = 0$.

The specific heat is the heat required to raise the temperature of a unit mass by one degree, and can be defined at constant volume or constant pressure:

$$C_V = \left(\frac{\delta q}{dT} \right)_v = \left(\frac{\partial u}{\partial T} \right)_v \quad (4.23)$$

$$C_P = \left(\frac{\delta q}{dT} \right)_P = \left(\frac{\partial u}{\partial T} \right)_P - \frac{P}{\rho^2} \left(\frac{\partial \rho}{\partial T} \right)_P \quad (4.24)$$

The ratio of specific heats is denoted by γ :

$$\gamma \equiv \frac{C_P}{C_V} = 1 + \frac{P}{\rho T C_V} \frac{\chi_T^2}{\chi_\rho} \quad (4.25)$$

and is equal to $\gamma = 5/3$ for an ideal gas (see Pols Ch 3.4.1 for details).

We now consider the response of a system to adiabatic changes (i.e., when $\delta q = 0$).

The adiabatic exponent measures the response of the pressure to an adiabatic compression or expansion:

$$\gamma_{\text{ad}} \equiv \left(\frac{\partial \ln P}{\partial \ln \rho} \right)_{\text{ad}} \quad (4.26)$$

For an adiabatic process we have:

$$du = \frac{P}{\rho^2} d\rho \quad (4.27)$$

we saw earlier in this lecture that $U \propto P$ in both the NR and UR limits. We can therefore write:

$$u = \phi \frac{P}{\rho} \quad (4.28)$$

where ϕ is a constant. If we differentiate the above equation and insert into Eq. 4.27 we arrive at:

$$\frac{dP}{P} = \frac{\phi + 1}{\phi} \frac{d\rho}{\rho} \quad (4.29)$$

which implies:

$$\gamma_{\text{ad}} = \frac{\phi + 1}{\phi} \quad (4.30)$$

For non-relativistic particles, including a classical ideal gas and NR degenerate electrons we have $\phi = 3/2$ and hence $\gamma_{\text{ad}} = 5/3$. For UR particles (photons or UR degenerate matter) we have $\phi = 3$ and therefore $\gamma_{\text{ad}} = 4/3$. The adiabatic exponent is important for dynamical stability of stars, as we will see in Lecture 6.

For a more general EOS (Eq. 4.20) we have:

$$\gamma_{\text{ad}} = \gamma \chi_{\rho} \quad (4.31)$$

In the case of an ideal gas, where $\chi_{\rho} = 1$, we have $\gamma = \gamma_{\text{ad}}$.

The adiabatic temperature gradient is defined as:

$$\nabla_{\text{ad}} = \left(\frac{\partial \ln T}{\partial \ln P} \right)_{\text{ad}} \quad (4.32)$$

so that $T \propto P^{\nabla_{\text{ad}}}$ if ∇_{ad} is constant. As we will see in Lecture 7, this quantity is important for considering stability against convection. In the case of an idea gas we have:

$$\nabla_{\text{ad}} = \frac{\gamma_{\text{ad}} - 1}{\gamma_{\text{ad}}} \quad (4.33)$$

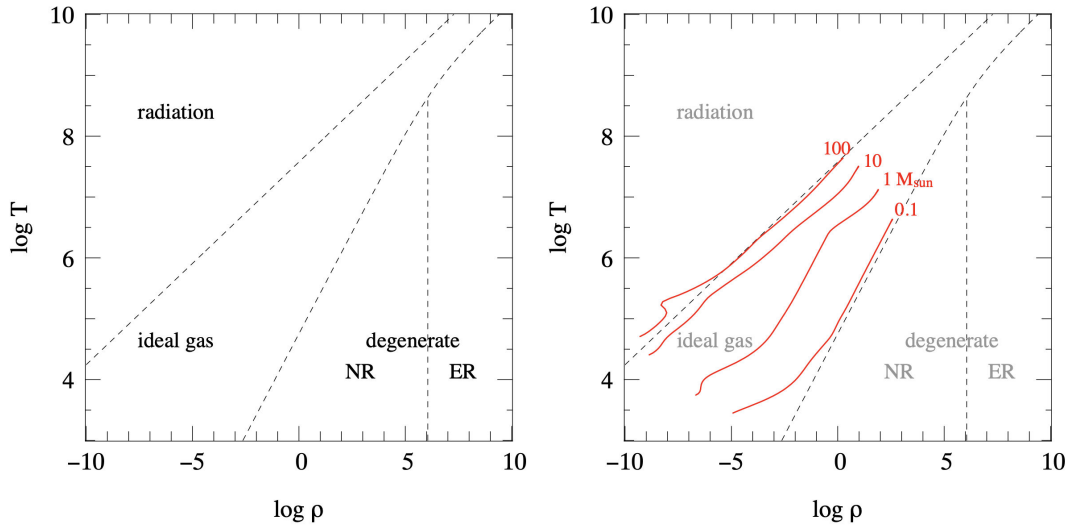


Figure 4.2: EOS for a gas of free particles. The left panel shows the regimes in which certain EOS's dominate. The right panel shows interior profiles for main sequence stars of various masses. ER refers to extremely relativistic and is equivalent to “UR” in these notes (from Pols 2011, Fig 3.4).

Chapter 5

Opacity & Nuclear Energy Generation

In this lecture we continue the discussion of the microphysics necessary for stellar model construction by introducing the most important sources of opacity and the basic physics of nuclear energy production.

Recall from Lecture 3 that **opacity** quantifies the probability of light interacting with matter and is often written as κ with units of $\text{cm}^2 \text{ g}^{-1}$. Opacity is sometimes expressed in terms of a cross section, σ , which is related to κ via: $\kappa \rho = n\sigma$ where n is the number density and ρ the mass density of the medium. We're going to briefly describe several standard sources and parameterizations of the opacity. The state of the art consists of tabulated opacity tables computed from consideration of the various forms of continuous and line opacity sources. Opacity tables are a function of the density, temperature, and composition.

At low densities and (relatively) high temperatures, a fully ionized gas has opacity that is well-characterized by **Thomson electron scattering**, which has a cross section:

$$\sigma_T = \frac{8\pi}{3} \left(\frac{e^2}{m_e c^2} \right)^2 \approx 6.65 \times 10^{-25} \text{ cm}^2 \quad (5.1)$$

note that $\frac{e^2}{m_e c^2}$ is the so-called classical radius of an electron. The Thomson cross section is therefore reasonably close to the intuitive geometric cross section of πr^2 . Using the above expressions for opacity combined with the relation between electron density and mass of $n_e = \frac{\rho(1+X)}{2m_H}$ we find: $\kappa_T = 0.2(1+X) \text{ cm}^2 \text{ g}^{-1}$. Notice that the Thomson opacity is independent of density and temperature and only depends on composition in the sense that different mixtures will have a different ratio of electrons to mass density. The Thomson opacity is used in many astronomical contexts, e.g., in simple models for massive stars, and in deriving the Eddington limit for accretion onto black holes, as we will see in later lectures.

The opacity due to free-free, bound-free, and bound-bound electronic transitions can be approximated with **Kramers' formula**. The formula for bound-free absorption is:

$$\kappa_{\text{bf}} \approx 4 \times 10^{25} (1 + X) (Z + 0.001) \rho T^{-3.5} \text{ cm}^2 \text{ g}^{-1} \quad (5.2)$$

where T is in K and ρ is in g cm^{-3} . The formulas for bound-bound and free-free have the same scaling with T and ρ but different numerical factors. See Pols (2011) Section 5.3.1 for a sketch of the “derivation”. This formula is applicable when the elements are at least partially ionized, i.e., for $T > 2 \times 10^4$ K. These relations are approximate and are useful for analytic estimates of the behavior of stellar interiors. They are not used in state-of-the-art models of stars. Instead, detailed models of atomic and molecular transitions are used, along with tabulations of bound-free and continuum opacities. See the accompanying slides for examples.

A particularly important source of opacity at $3000 \text{ K} \lesssim T \lesssim 10^4 \text{ K}$ is bound-free absorption from the H^- ion. The ion is very fragile and so is easily ionized at a few thousand degrees. The source of the extra electron is mostly from singly ionized metals such as Na, K, and Ca. The resulting opacity therefore is sensitive to metallicity as well as temperature. A very approximate formula for this source of opacity is:

$$\kappa \propto Z \rho^{1/2} T^{7.7} \quad (5.3)$$

We will see the critical role that this source of opacity plays in governing the effective temperature of fully convective stars in Lecture 9.

Our understanding of how (and for how long) the Sun shines has a rich history. As alluded to in Lecture 1, in the mid-nineteenth century Lord Kelvin argued forcefully based on the energy available from gravitational contraction that the Sun was no older than 3×10^7 yr. Darwin and others argued on geological grounds that the Earth was much older than this. The resolution to this problem required multiple fundamental advances and discoveries in physics, including 1) the discovery of radioactivity by Rontgen, Curie, and Rutherford at the turn of the century; 2) Einstein's discovery in 1905 of the equivalence of mass and energy via $E = mc^2$; 3) the realization in 1928 by Gamow that quantum tunnelling was necessary for fusion to take place within stars; 4) the comprehensive calculations by Bethe in 1938 regarding the basic nuclear processes by which hydrogen is burned into helium. This history offers a valuable lesson in how the solution to a problem sometimes requires multiple fundamental breakthroughs on a variety of topics. We will now briefly review the essential physical processes responsible for **nuclear energy production** in stars.

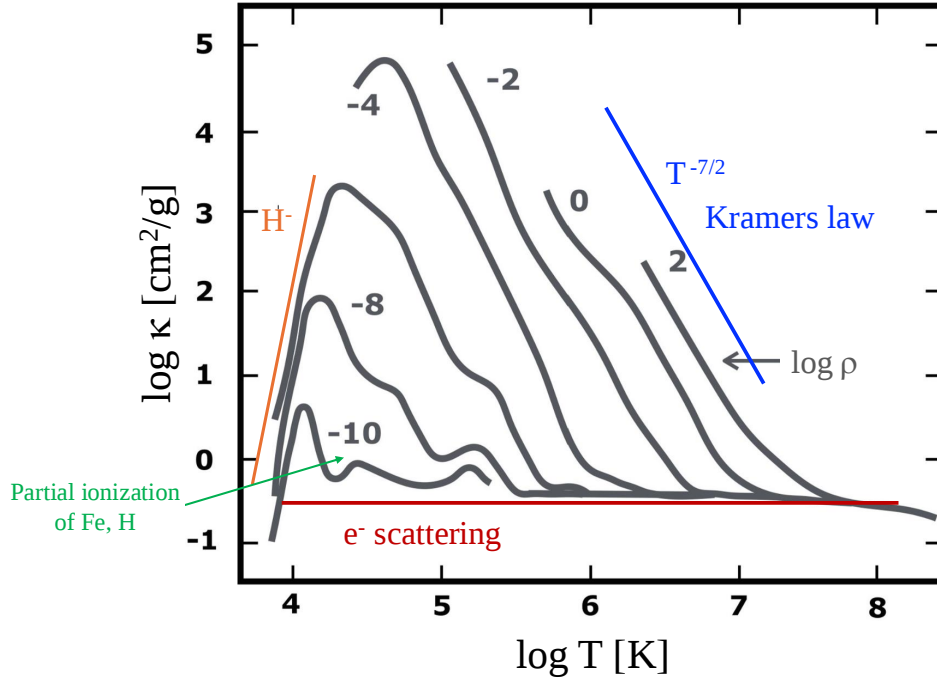


Figure 5.1: Opacity as a function of temperature and density for a solar element abundance distribution (grey lines). The three approximations discussed in the text are shown (Thomson, Kramers, and H^- opacity). Adapted from Lamers & Levesque.

The **binding energy** (E_b) is the energy released (or absorbed) when a nucleus is assembled from its constituent nucleons: $E_b = \Delta m c^2$ where $\Delta m = \sum m_j - m_{\text{nuc}}$.

If we define A as the number of nucleons, N as the number of neutrons and Z_p as the number of protons (the subscript is meant to help distinguish our use of Z for metallicity), then $A = N + Z_p$. An important concept is the binding energy per nucleon: E_b/A , which ranges from $\approx 1 - 8.5$ MeV.

Notation: eV is neither an SI nor cgs unit. The cgs energy unit is erg, which is equal to 6.24×10^{11} eV.

A plot of E_b/A vs. A has several important features (see Fig. 5.2), but the most important is that the relation has a peak at iron (often referred to as the “iron peak”). The consequence of this peak is that for $A < 56$ fusion releases energy, while for $A > 56$ fusion consumes energy. This fact has significant implications for the evolution of stars, as we will see in later lectures.

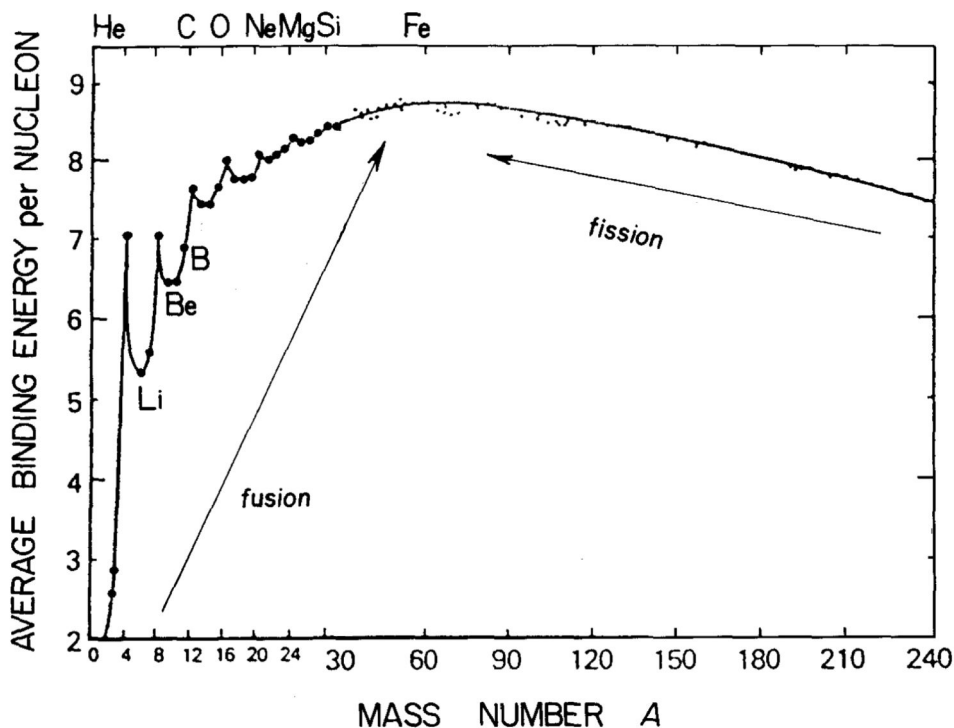
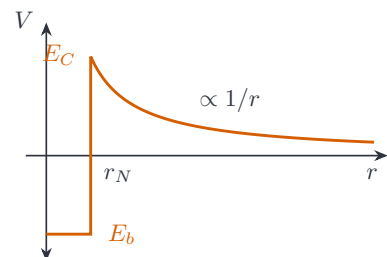


Figure 5.2: Average binding energy in MeV per nucleon as a function of atomic mass. From Rolfs & Rodney (1988).

Fusion requires collisions between nuclei, which means that colliding nuclei must overcome the Coulomb barrier, defined as $E_C = 1.4 \text{ MeV } Z_1 Z_2 / r_{1\text{fm}}$ where Z_1 and Z_2 are the proton numbers of the two nuclei and $r_{1\text{fm}}$ is the separation in units of $\text{fm} = 10^{-13} \text{ cm}$. The “size” of a nucleus is $r_N \approx 1.3 \text{ fm } A^{1/3}$. See the diagram at right.



Classically, two particles would require a relative energy in the MeV range in order to overcome the Coulomb barrier and fuse. However, in the center of the Sun particles have a typical energy of $kT \sim 1 \text{ keV}$. Appealing to the Maxwellian distribution of energies, the fraction of nuclei with thermal energy in the MeV range is $\sim e^{-1000}$ - that is a very small number! The solution to this problem had to await for developments in quantum mechanics, specifically, quantum tunneling.

In quantum mechanics, waves can tunnel through a potential barrier. The wave function decays exponentially within the classically-forbidden region. The result (first shown by Gamow) is that the probability of tunneling scales as $P_t(E) \approx e^{-\sqrt{E_G/E}}$ where E_G is known

as the Gamow Energy and is defined as $E_G \equiv (\pi\alpha Z_1 Z_2)^2 2m_r c^2$ where $\alpha \approx 1/137$ is the fine structure constant and $m_r = m_1 m_2 / (m_1 + m_2)$ is the reduced mass. Plugging in the constants: $E_G \approx 1 \text{ MeV } Z_1^2 Z_2^2 (m_r/m_H)$. We can now express the probability of tunneling through the Coulomb barrier in terms of the properties of the two nuclei. The next step is to determine the *rate* of such interactions.

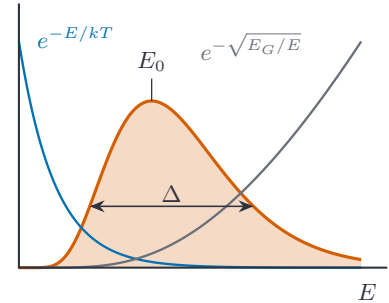
Two-body **reaction rates** generally have the following functional form: $R_{12} = n_1 n_2 \langle \sigma v \rangle$ with units $\text{s}^{-1} \text{ cm}^{-3}$, or rate per volume. n_i are number densities of the two species under consideration, v is the relative collision velocity, and σ is the cross section for interaction. Notice that reaction rates scale as n^2 and v , so higher central densities will produce much higher reaction rates, as will higher temperatures (via the effect of temperature on velocity). The average is over the velocity distribution: $\langle \sigma v \rangle = \int \sigma v P(v) dv$ where $P(v)$ is usually the Maxwellian distribution of velocities in the case of a thermal distribution of particles.

The cross section of interaction is approximately: $\sigma \sim \pi \lambda^2 P_t(E)$ where λ is the de Broglie wavelength and is equal to h/p ; recall that $p^2 = 2m_r E$. The cross section is therefore: $\sigma(E) = S(E) E^{-1} e^{-\sqrt{E_G/E}}$ where $S(E)$ represents the intrinsically nuclear parts of the reaction probability. It is normally a slow function of E except in cases of resonant reactions, when $S(E)$ changes rapidly with E .

Recalling that the Maxwellian distribution is $P(v)dv \propto e^{-mv^2/kT} v^2 dv$ we can write the general behavior of the reaction rate as:

$$R_{12} \propto n_1 n_2 \int S(E) e^{-E/kT - \sqrt{E_G/E}} dE \quad (5.4)$$

The factor within the exponential is the key piece, as indicated in the diagram at right. The basic idea is that the Maxwellian tail of the energy distribution suppresses the distribution at high energies, but the tunneling probability



increases with increasing energy. The result is that fusion will occur within a narrow window with a peak at $E_0 = (E_G (kT)^2/4)^{1/3}$ and an approximate width of $\Delta = 2 E_G^{1/6} (kT)^{5/6}$. We can estimate the temperature dependence of the reaction rate by approximating the term within the integral as a Gaussian with peak E_0 , width Δ and $S(E) \approx S(E_0)$ (see Pols, Eqns 6.24–6.30 for details). Integrating this equation leads to:

$$\langle \sigma v \rangle \propto T^{-2/3} e^{-CT^{-1/3}} \quad (5.5)$$

where C is a constant that depends on $Z_1 Z_2$, i.e., on the height of the Coulomb barrier. If

we consider a small range of temperatures around a reference value T_0 then we have:

$$\langle \sigma v \rangle \propto (T/T_0)^\nu \quad \text{where} \quad \nu \equiv \frac{\partial \log \langle \sigma v \rangle}{\partial \log T} = \frac{3E_0/kT - 2}{3} \quad (5.6)$$

It is clear that the temperature dependence of nuclear reaction rates is not a power-law, as ν itself depends on T . However, in practice any particular nuclear reaction is important over a relatively limited range of temperatures, so one often associates approximate power-law indices for specific reaction networks. The key conclusion from the above equation (which is not obvious until one plugs in example numbers) is that ν is always large, i.e., the temperature-dependence is steep.

Lets consider two examples reactions:

$$p + d \rightarrow {}^3\text{He} + \gamma, \quad E_G = 0.657 \text{ MeV} \quad (5.7)$$

$$p + {}^{12}\text{C} \rightarrow {}^{13}\text{N} + \gamma, \quad E_G = 35.5 \text{ MeV} \quad (5.8)$$

The first reaction is part of the pp chain, while the second is part of the CNO cycle. At $T = 2 \times 10^7$ K these two reaction rates scale as T^4 and T^{17} , respectively. That is a huge difference!

Lets now consider several of the most important nuclear reaction networks for stars.

For **Hydrogen burning**, the first important network is the **pp chain**, of which there are several branches, but the main set is the following:

$$p + p \rightarrow d + e^+ + \nu_e \quad (5.9)$$

$$p + d \rightarrow {}^3\text{He} + \gamma \quad (5.10)$$

$${}^3\text{He} + {}^3\text{He} \rightarrow {}^4\text{He} + 2p \quad (5.11)$$

where d is deuterium, e^+ is the positron, and ν_e is the electron neutrino. The net effect of the pp chain is: $4p \rightarrow {}^4\text{He} + 2e^+ + 2\nu_e + \gamma$ with $Q = 25$ MeV of energy release (corresponding to a release in fractional rest-mass energy of $\Delta m/m = 0.007 = 0.7\%$). Note that the cross-section for neutrino interaction is very small, so the neutrino particles immediately free-stream away. These particles therefore carry valuable information regarding the nuclear reactions in the core of the Sun.

The first reaction is very slow: $R \approx 5 \times 10^7 \text{ s}^{-1} \text{ cm}^{-3}$. On average, it takes a proton $> 10^9$ yr to fuse with another proton in the center of the Sun. This explains the long lifetime of the Sun. If we combine the rate with the energy per reaction we arrive at an energy production rate in the core of the Sun of $\sim 120 \text{ W m}^{-3}$, or about one incandescent light bulb per m^3 !

The second major hydrogen burning network is the **CNO cycle**. The key concept is that the elements CNO are used as *catalysts* to convert protons into helium. The reaction chains are longer than the pp chain and contain numerous branches (including so-called hot and cold CNO burning), so I won't write them down here. The main points are i) 24 MeV of energy is released (nearly identical to the pp chain, which is not surprising), ii) the sum of C+N+O is not altered; iii) the individual abundances of C, N, and O *are* altered to reflect equilibrium conditions (and are set by the inverse of the reaction rate governing each species); iv) the reaction rate is a *very* strong function of temperature (T^{17} , as we saw above).

For the Sun and lower-mass stars, hydrogen is fused into helium via the pp chain. For stars more massive than $\approx 1.2M_{\odot}$ the CNO cycle is the dominant hydrogen burning channel.

Question: How do metal-free (so-called Pop III) massive stars burn hydrogen? What implications might this have for the properties of such stars? You will explore this topic using MESA in HW2.

Helium burning occurs via the **triple- α process** (a helium nucleus is also known as an α particle). The reaction chain is:

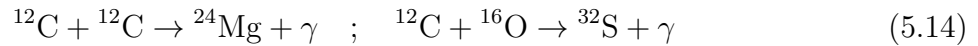


where $Q = 7.3$ MeV. The lifetime of ${}^8\text{Be}$ is very small: 10^{-16} s, which means that this reaction network is essentially a three-body interaction. As a consequence, the effective reaction rate scales as $R \propto \rho^3 T^{40}$ at 10^8 K. Notice the extremely steep dependence on temperature.

Recall from Lecture 1 that we can estimate the timescale to burn through a fuel source as $t = E/\dot{E}$. There is roughly $3\times$ more energy produced in hydrogen compared to helium burning (25 vs. 7 MeV). Moreover, as we will see later, the luminosity of helium burning is $50\times$ higher than hydrogen burning, so the lifetime of the helium burning phase is only 1% of the hydrogen burning phase (for a $1M_{\odot}$ star - these numbers all vary with mass and metallicity).

Toward the end of helium burning, when the helium abundance is low and the carbon abundance is high, another reaction becomes important: ${}^{12}\text{C} + {}^4\text{He} \rightarrow {}^{16}\text{O}$. For stars $\lesssim 8M_{\odot}$, this is the end of the road (at least in their cores). Helium burning produces a CO core that, after some additional burning in shells surrounding the core, evolves into an inert white dwarf.

For more massive stars, burning will proceed up to (and beyond) ^{56}Fe . At $T \sim 10^9$ K the following reactions take place:



lots of other reactions are also occurring, though they are not the dominant energy sources.

At $T > 2 \times 10^9$ K ($kT \approx 0.4$ MeV), photons can destroy heavy nuclei. For example, at these temperatures photons can destroy the relatively fragile ^{20}Ne atom via: $^{20}\text{Ne} + \gamma \rightarrow ^{16}\text{O} + ^4\text{He}$. This is an example of a process called **photodisintegration**, which will create a soup of p , n , α , and some light nuclei. Subsequent α capture reactions create lots of elements including $^{56}\text{Ni} \rightarrow ^{56}\text{Co} \rightarrow ^{56}\text{Fe}$. This process of element formation is governed by **nuclear statistical equilibrium (NSE)**, which is similar to the Saha equation for nuclei.

The formation of very heavy elements ($Z_p > 30$) occurs via **neutron capture**. Neutrons are easily captured by nuclei because neutrons have no charge, so they do not experience Coulomb repulsion. In this way very neutron-rich nuclei can form. But such nuclei are unstable, so the details of what stable elements form depends on whether the neutron flux is “slow” or “fast” compared to the decay time. The slow channel is known as the **s-process**. The archetype site for the s-process is the envelopes of AGB stars. Examples of elements formed mostly in this way are Zr, Sr, and Ba.

The fast channel is known as the **r-process**. For a long time the two proposed sites for this process included supernovae and merging neutron stars. Observations in 2017 of the LIGO source GW170817 provided the first strong evidence for a link between merging neutron stars and substantial production of r-process material. Another indirect argument in favor of the merging neutron star channel is based on the expectation that, compared to supernovae, merging neutron stars are rare, but each event produces a larger quantity of r-process elements. This means that one would expect a larger *variation* in r-process abundances if merging neutron stars are the main channel. Simple models have been constructed to explain the abundance patterns in dwarf galaxies and in metal-poor stars in our Galaxy, and the general conclusion is that the data favor a channel in which rare events produce large quantities of r-process elements, but there is likely at least some contribution from supernovae. This is an active area of research.

Question: Imagine two hypothetical sources of the r-process element Eu. The first source is rare but produces large quantities of Eu per event. The second source is common and produces small quantities per event. Averaged over the age of the universe they produce the same total quantity of Eu. What observations might be able to distinguish between these two sources?

The concept of **neutrino cooling** will arise in several situations as we discuss stellar evolution in later lectures. At very high densities and temperatures, the coupling of electromagnetic and weak interactions means that there is a very small, but non-zero probability that a process which normally produces a photon instead produces a neutrino-antineutrino pair. This probability is roughly:

$$\frac{P(\nu\bar{\nu})}{P(\gamma)} \approx 3 \times 10^{-18} \left(\frac{E_\nu}{m_e c^2} \right)^4 \quad (5.15)$$

where E_ν is the neutrino energy. Because the cross section for neutrino interactions is so small, these neutrinos represent a direct loss of energy from the stellar interior. This process is called neutrino cooling.

There are two important situations in which neutrino cooling arises in stellar interiors. The first is photo-neutrinos. In this case, the normal photon-electron scattering is replaced with a neutrino, i.e., $\gamma + e^- \rightarrow e^- + \nu + \bar{\nu}$. The average neutrino energy is $E_\nu \sim kT$, and so the probability scales as T^4 . Moreover, the overall rate of neutrino emission is proportional to the number density of photons, which scales as T^3 , so the overall photo-neutrino cooling rate scales as T^7 . Obviously, this cooling will be important at high temperatures.

The second important case is pair annihilation neutrinos. At $T > 10^9$ K, the rest energy of electrons is comparable to the photon energy, and photons can undergo electron pair creation: $\gamma + \gamma \leftrightarrow e^- + e^+$. Because of the same coupling mentioned above, very rarely the following reaction occurs instead: $e^- + e^+ \leftrightarrow \nu + \bar{\nu}$, which again results in a one-way loss of energy. This is an important cooling mechanism for the late-stages of massive stars.

Summary: We have concluded our brief tour of the “microphysics” necessary for constructing stellar models. The EOS, opacities, and nuclear reactions provide the connection between the large-scale structure of stars and the detailed properties of atoms; it is primarily through these three ingredients that the structure of the atom and the nature of quantum mechanics is able to manifest at the macroscopic scales. These ingredients are also the vehicle by which *uncertainties* in our understanding of energy levels in atoms and the nature of matter under extreme conditions impact our ability to make robust predictions of stellar structure and evolution.

There are several areas of current research in this field. The opacity at cool temperatures is difficult to compute owing to the large number of molecules present. In practice, opacity tables are constructed over limited ranges of temperature and density, and tables from different

sources must be blended together, which can cause numerical artefacts in stellar evolution programs. In certain phases of stellar evolution the CNO abundance changes substantially in the stellar interior, and in such cases opacity tables must be specially constructed to span the full range of abundances. Finally, the cross sections for numerous key nuclear reactions are uncertain at the level where stellar evolution predictions are affected (see Fields et al. 2018 for details).

Further Reading: Lamers & Levesque Ch 5, 8; for a more advanced treatment of nuclear reactions, see Clayton's textbook "Principles of Stellar Evolution and Nucleosynthesis"

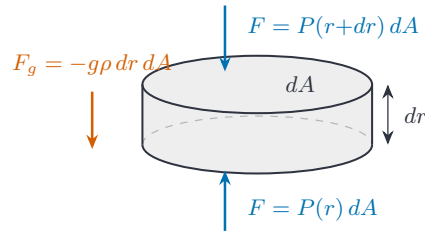
Chapter 6

Mechanical Structure & Energy Conservation

We will now turn our attention to the stellar structure equations, which are the key governing equations for the structure and evolution of stars. In this lecture we will consider the mechanical structure of stars arising from hydrostatic equilibrium. We will derive and analyze three of the four key equations of stellar structure, and we will introduce the concept of a polytrope. We will also sketch the derivation of the virial theorem.

Let's begin by considering the implications of **hydrostatic equilibrium**, in which a star is supported against gravity by pressure gradients.

The diagram at right shows the basic geometry. We consider an infinitesimal cylinder, in which there is an upward pressure on the bottom and a downward pressure on the top. The net pressure force is $F_P = P(r)dA - P(r+dr)dA$. Equilibrium would imply that the force of gravity balances this net pressure force, i.e., $F_p = F_g$, where $F_g = -g\rho dr dA$. Setting these two equal yields:



$$\frac{P(r+dr) - P(r)}{dr} = \frac{dP}{dr} = -g\rho \quad (6.1)$$

where $g = Gm(r)/r^2$ and $m(r) = m_r = m(< r)$. Therefore, our first equation of stellar structure is:

$$\#1 : \boxed{\frac{dP}{dr} = -\frac{G m(r) \rho(r)}{r^2}} \quad (6.2)$$

Question: Use order-of-magnitude approximations to estimate the central pressure of a star. Assuming an ideal gas EOS, estimate the central temperature. Use these expressions to estimate the central temperature of the Sun. For comparison, modern estimates of the central temperature are $\approx 15 \times 10^6$ K.

The second equation of stellar structure is the relation between mass and radial element and is known as the mass continuity equation:

$$\#2 : \quad \boxed{\frac{dm}{dr} = 4\pi r^2 \rho(r)} \quad (6.3)$$

note that this equation allows one to recast the other stellar structure equations in mass coordinates, e.g., you will sometimes see the hydrostatic equilibrium equation written as:

$$\frac{dP}{dm} = -\frac{Gm(r)}{4\pi r^4} \quad (6.4)$$

These two equations contain three unknowns; P, ρ, m . If we invoke the ideal gas law, $P = nkT$, then we need to know the temperature. This leads us to consider constraints imposed by energy conservation and energy transport (see below and Lecture 7). Another approach is to consider a **polytropic EOS**: $P = K\rho^\gamma$, and $\gamma = (n+1)/n$ where n is the polytropic index and K is a constant (unrelated to the kinetic energy, which is also represented by K later in these notes). This was a popular approach to generate simple stellar models before the age of computers. You will see in HW2 that real stellar models are often surprisingly close to polytropes, which justifies a detailed consideration of their properties. Polytropic models can provide valuable insight into the structure of stars without the need for computers.

We have already encountered two examples of polytropes: degenerate EOS for non-relativistic ($P \propto \rho^{5/3}$) and ultra-relativistic matter ($P \propto \rho^{4/3}$). As we will see in the following lecture, convective regions are very nearly adiabatic, which means that their EOS is also well-described by a polytrope.

With a polytropic EOS we can now solve the two equations of stellar structure as follows. Let's start by re-casting the hydrostatic equilibrium equation in terms of the gravitational potential, Φ : $dP/dr = -\rho d\Phi/dr$, and after inserting the polytropic EOS: $\frac{K}{\rho} \frac{d\rho}{dr} = -\frac{d\Phi}{dr}$. We will now appeal to Poisson's equation: $\nabla^2\Phi = 4\pi G\rho$, which in spherical coordinates can be expressed as: $\frac{1}{r^2} \frac{d}{dr} (r^2 \frac{d\Phi}{dr}) = 4\pi G\rho$. Combining:

$$\left(\frac{n+1}{n} \right) \frac{K}{r^2} \frac{d}{dr} \left[r^2 \rho^{1/n-1} \frac{d\rho}{dr} \right] = -4\pi G\rho \quad (6.5)$$

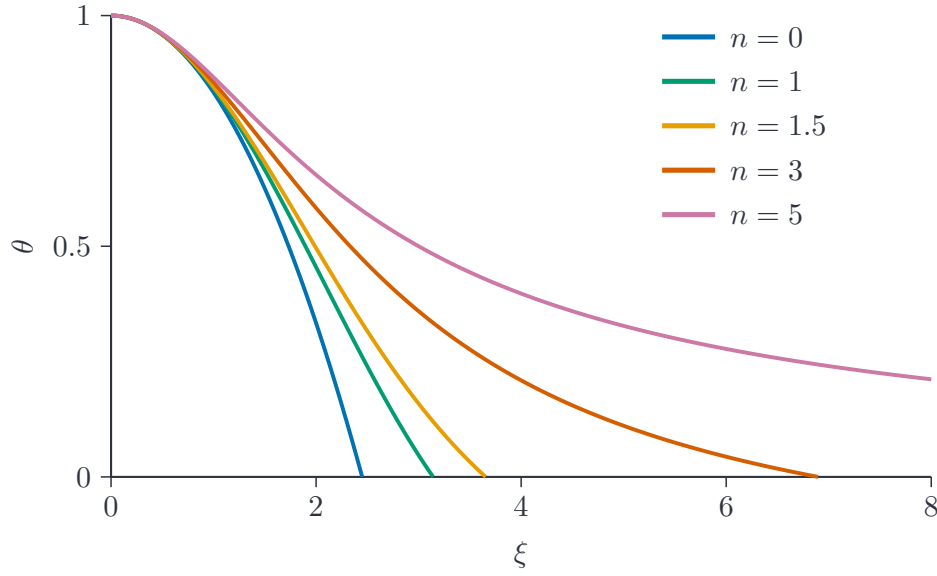


Figure 6.1: Solutions of the Lane-Emden equation for five polytropic indices, obtained by numerical integration. Recall that $\rho(r) = \rho_c \theta^n$ and $r = \lambda_n \xi$, so θ falls from unity at the center to zero at the surface, which is reached at $\xi = \xi_1$. Larger n corresponds to a softer EOS and a more centrally concentrated star: $\xi_1 = 2.449$ for $n = 0$, 3.654 for $n = 1.5$, and 6.897 for $n = 3$. The $n = 5$ solution never reaches zero, which is why polytropes with $n \geq 5$ have infinite radius.

where $n \equiv 1/(\gamma - 1)$ is the polytropic index.

Now let $\rho(r) = \rho_c \theta^n$ and $r = \lambda_n \xi$. The resulting expression is known as the **Lane-Emden Equation**:

$$\frac{1}{\xi^2} \frac{d}{d\xi} \left[\xi^2 \frac{d\theta}{d\xi} \right] = -\theta^n \quad (6.6)$$

The solutions are analytic for $n = 0, 1, 5$. Figure 6.1 shows solutions for several values of n .

Some common examples of polytropes:

- $n = 0$: these are constant density models, i.e., “incompressible” (recall the refactoring of the L-E Eqn: $\rho(r) = \rho_c \theta^n$). Note that $n = 0$ implies $\gamma = \infty$. Examples include rocky planets.
- $n = 1.5$: these have $\gamma = 5/3$ and so represent non-relativistic degenerate matter and also convective regions in stars.
- $n = 3$: these have $\gamma = 4/3$ and so represent ultra-relativistic degenerate matter. Eddington’s model for stars (see below) in which $P_{\text{gas}}/P = \text{constant}$ is also an $n = 3$ polytrope, as is a star dominated by radiation pressure (where $\gamma = 4/3$).

- $n \geq 5$: polytropes of this kind have infinite radius.
- $n = \infty$: corresponds to $\gamma = 1$, which is an isothermal distribution, i.e., $T = \text{constant}$, (recall that for an ideal gas we have $d \ln T / d \ln P = \frac{\gamma-1}{\gamma}$). We often encounter isothermal cores in various phases of stellar evolution; these are well-described by polytropes with $n = \infty$.

Let's return to the definition of a polytrope: $P = K \rho^\gamma = \rho^{1+1/n}$. We can derive scaling relations between the total mass and radius by inserting $\rho \sim M/R^3$ and $P \sim GM^2/R^4$ in the polytropic relation. Combining, we find:

$$R^{\frac{3-n}{n}} M^{\frac{n-1}{n}} \propto \frac{K}{G} \quad (6.7)$$

Actually solving the L-E equation produces an additional function that depends on n on the right-hand side of the equation.

Consider two examples: $n = 1.5$ implies $RM^{1/3} = \text{constant}$. This means that the radius *decreases* as the mass increases! This occurs for white dwarfs (what happens what $R \rightarrow 0$, I wonder?). $n = 3$ implies a constant mass, equal to $(K/0.36G)^{1.5}$, where 0.36 comes from actually integrating the L-E equation. This is obviously a very special case. As neutron stars are described well by $n = 3$ polytropes, we can anticipate 1) the existence of a maximum neutron star mass, 2) the maximum mass will provide insight into the neutron star EOS (e.g., via the constant K , though state-of-the-art models are of course more complex, as we saw in Lecture 4).

Finally, we can compute the gravitational potential energy of a polytrope. The result is:

$$U = -\frac{3}{5-n} \frac{GM^2}{R} \quad (6.8)$$

We will now briefly consider the conditions required for global **dynamical stability** of stars. Consider a star in hydrostatic equilibrium with mass M , density ρ , and radius R . Suppose the star is compressed on a short timescale ($\tau \ll \tau_{\text{KH}}$) so that the compression is adiabatic. Assume that the star is compressed from its original radius R to a new radius R' so that the density increases to: $\rho' = \rho(R'/R)^{-3}$. Because the compression is adiabatic the new pressure is

$$\frac{P'}{P} = \left(\frac{\rho'}{\rho} \right)^{\gamma_{\text{ad}}} \quad (6.9)$$

The pressure required to maintain hydrostatic equilibrium is $P = GM^2/R^4$, which in our case implies:

$$\left(\frac{P'}{P}\right)_{\text{HE}} = \left(\frac{R'}{R}\right)^{-4} = \left(\frac{\rho'}{\rho}\right)^{4/3} \quad (6.10)$$

If $\gamma_{\text{ad}} > \frac{4}{3}$ then the compression results in an excess of pressure compared to what is needed for hydrostatic support, and the result is that the star will re-expand (on the dynamical timescale) so that hydrostatic equilibrium is restored. However, if $\gamma_{\text{ad}} < \frac{4}{3}$ then the increase in pressure from the compression is insufficient to restore hydrostatic balance, and the star will continue to compress (again on the dynamical timescale). This leads to an important condition for dynamical stability: $\gamma_{\text{ad}} > \frac{4}{3}$.

Our sketch here has been very simple. In reality we often find only certain regions of a star to have $\gamma_{\text{ad}} < \frac{4}{3}$ and then it is less clear exactly what happens. It turns out that global dynamical instability is obtained when:

$$\int_0^M \left(\gamma_{\text{ad}} - \frac{4}{3}\right) \frac{P}{\rho} dm < 0 \quad (6.11)$$

This implies that if $\gamma_{\text{ad}} < \frac{4}{3}$ in the core, where P/ρ is high, then the star will be unstable, but if $\gamma_{\text{ad}} < \frac{4}{3}$ in the outer layers, where P/ρ is small, then the star as a whole may still be stable.

As we have seen above, we have $\gamma_{\text{ad}} = 5/3$ for an ideal gas, non-relativistic degenerate matter, as well as convective regions. We encounter regions where $\gamma_{\text{ad}} = 4/3$ when relativistic generate matter or radiation dominates the total pressure. In later lectures we will encounter additional effects that can result in $\gamma_{\text{ad}} < 4/3$ at certain layers in stars: pair creation and photo-disintegration of nuclei (relevant in stellar cores), and partial ionization zones (relevant in the envelopes of stars). In such cases we can expect the star to be much more susceptible to either local or global dynamical instabilities.

The **virial theorem** has many applications in astronomy. We will employ it throughout this course, so a brief derivation is provided here based on the hydrostatic equilibrium equation (stellar structure equation #1 above). There are many routes to deriving the virial theorem, this being only one. By starting with HE we have already implicitly assumed that all time derivatives are zero.

We begin by multiplying both sides of the HE equation by $4\pi r^3 dr$ and integrating, which leads to:

$$\int 4\pi r^3 dP = - \int \rho \frac{Gm(r)}{r^2} 4\pi r^3 dr \quad (6.12)$$

using stellar structure equation #2 above ($dm = 4\pi r^2 \rho dr$) leads to

$$\int 4\pi r^3 dP = - \int \frac{Gm(r)}{r} dm = U \quad (6.13)$$

where the last term is just the gravitational binding energy. The first term can be simplified by appealing to a vector calculus identity:

$$\int 4\pi r^3 dr \frac{dP}{dr} = 4\pi r^3 P \Big|_0^R - 3 \left[\int P 4\pi r^2 dr \right] = 4\pi r^3 P(R) - 3 \langle P \rangle V \quad (6.14)$$

note that in many astronomical applications the external pressure (the pressure at the boundary), $P(R)$, is zero. We assume this below.

Question: Earlier in this lecture you estimated the central pressure of the Sun. Compare this value to the surface pressure (at $\tau = 1$) that you estimated in HW1. Is our usual assumption that $P(R) \approx 0$ when applying the virial theorem to stars a good one?

The virial theorem therefore states: $\langle P \rangle V = -\frac{1}{3}U$. For an ideal gas we also know that $PV = NkT$ and $K = \frac{3}{2}NkT$ where K is the kinetic energy, so $K = \frac{3}{2}PV$. Defining E as the total energy, $E = U + K$, yields the familiar forms of the virial theorem:

$$E = -K = \frac{1}{2}U \quad (6.15)$$

We can also derive a relativistic version of the virial theorem by recalling that $\langle P \rangle = \frac{1}{3}u_{\text{kin}} = \frac{1}{3} \frac{K}{V}$ for ultra-relativistic particles (e.g., photons). Equating this to the earlier expression of the virial theorem: $\langle P \rangle V = -\frac{1}{3}U$, results in the following expression: $K = -U$, appropriate for relativistic particles. Note that in this case the total energy $E = U + K = 0$! In this case small perturbations can lead to instability and the destruction of the star. This is another way of understanding the dynamical stability limit discussed earlier in this lecture.

We may also ask under what conditions are required for **thermal stability**. Because in general we have $t_{\text{dyn}} \ll t_{\text{KH}} \ll t_{\text{nuc}}$, we can assume that the star is in dynamical equilibrium here. Lets further assume that the star is supported by ideal gas pressure.

The virial theorem tells us that $E = -K = -\frac{3}{2}Nk\bar{T}$ where $\bar{T}$ is the mass-weighted temperature. Energy conservation implies:

$$L_{\text{nuc}} - L = \frac{dE}{dt} = -\frac{3}{2}Nk \frac{d\bar{T}}{dt} \quad (6.16)$$

where L_{nuc} is energy provided by nuclear fusion (i.e., energy production) and L is the total luminosity (i.e., the energy loss). In thermal equilibrium we have $L_{\text{nuc}} = L$ and then $d\bar{T}/dt = 0$. To investigate the stability of thermal equilibrium we consider small perturbations where $L_{\text{nuc}} - L = \delta L \neq 0$. In this case the mean temperature changes according to:

$$\frac{d\bar{T}}{dt} = -\frac{2}{3} \frac{1}{kN} \delta L \quad (6.17)$$

Now imagine that a perturbation results in $\delta L > 0$, e.g., because of a local increase in temperature which results in an increase in L_{nuc} . According to Eq. 6.17, this perturbation will result in a decrease in the temperature. This in turn will result in a decrease in L_{nuc} , because of the strong temperature sensitivity of nuclear reaction rates. In contrast, if $\delta L < 0$, the temperature will rise and L_{nuc} will increase. This is precisely the condition for thermal stability.

The behavior here is counter to our normal experience - usually when energy is added to a system it heats up. However, systems that are self-gravitating have the peculiar property of having a negative specific heat, so that adding heat actually makes the system colder!

In degenerate matter the pressure is independent of the temperature. In this situation the star can be quite thermally unstable, as we will see in the case of the helium flash. We will encounter another case of thermal instability in Lecture 12: the thin shell instability.

We will now consider a simple but surprisingly accurate model for stellar interiors known as **Eddington's model**. The basic assumption is that the ratio of gas-to-total pressure is constant throughout the star: $\beta \equiv P_{\text{gas}}/P = \text{constant}$. We write the total pressure as the sum of (ideal) gas and radiation terms:

$$P = P_{\text{gas}} + P_r = \frac{k}{\mu m_{\text{H}}} \rho T + \frac{a}{3} T^4 \quad (6.18)$$

We can construct two relations by expressing the gas and radiation pressure in terms of β : $T = \frac{\mu m_{\text{H}}}{k} \frac{P_{\text{gas}}}{\rho} = \frac{\mu m_{\text{H}}}{k} \frac{\beta P}{\rho}$, and $T^4 = 3P_r/a = 3(1 - \beta)P/a$. The first equation reduces to $P \propto \rho T$, which we can use to reduce the second equation to $T^4 \propto \rho T$ which implies $\rho \propto T^3$. Plugging this back into our original expression for the total pressure leads us to

$$P = K \rho^{4/3} \quad , \quad K = \left[\frac{3}{a} \left(\frac{k}{\mu m_{\text{H}}} \right)^4 \frac{1 - \beta}{\beta} \right]^{4/3} \quad (6.19)$$

This is an $n = 3$ polytrope in which $M = (K/0.36G)^{1.5}$, which allows us to write:

$$\frac{M}{M_{\odot}} = \frac{18.1}{\mu^2} \frac{(1 - \beta)^{1/2}}{\beta^2} \quad (6.20)$$

Notice that $\beta = 0.5$ implies $M = 51M_\odot/\mu^2$. An important distinction between Eddington's model and the EOS for ultra-relativistic degenerate matter, which both have $n = 3$ is that the latter has a fixed value of K , while the former explicitly has $K = f(\beta)$.

Eddington used this model to calculate the structure of the Sun, and it works surprisingly well (see e.g., Lamers & Levesque Fig 11.2). As we will see in the HW, many stars are well-approximated by $\beta = \text{constant}$.

In order to solve the full stellar structure equations in a more general form we need to know the thermal energy content and the flow of energy within the star.

The third stellar structure equation is, like the second, mostly a matter of definition and is known as the **energy conservation** equation:

$$\#3 : \quad \boxed{\frac{dL_r}{dr} = 4\pi\rho(\epsilon_{\text{nuc}} - \epsilon_\nu + \epsilon_{\text{gr}})} \quad (6.21)$$

where ϵ_{nuc} is the nuclear energy generation per unit time per unit mass, ϵ_ν is the neutrino loss rate per unit mass, and ϵ_{gr} is the energy change due to contraction or expansion of the layer under consideration ($\epsilon_{\text{gr}} = -Tds/dt$). This last term is only relevant if the change occurs on a thermal timescale, otherwise $\epsilon_{\text{gr}} \approx 0$. All of these energy production/loss terms will depend on radius.

The fourth stellar evolution equation is related to energy transport, or the energy flux, F , defined through the equation: $L_r = 4\pi r^2 F$. Here we reach an important branching point that depends on the process responsible for energy transport. The three main possibilities are radiation, convection, and conduction. The latter is rarely important, except in the case of white dwarfs. In the next lecture we will discuss energy transfer via radiation and convection.

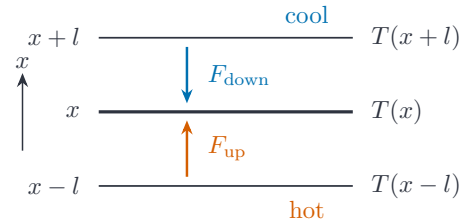
Summary: In this lecture we have sketched the basic mechanical structure of stars and were able to identify several analytic models assuming either a polytropic EOS or a simple (Eddington) model for the pressure. As you will see in HW2, these models are surprisingly accurate given their simplicity and provide important insight into the interior structure of stars. To make further progress we need to understand how energy is transported within stars. This step is much more complicated (and phenomenologically rich) and is the subject of the next lecture.

Chapter 7

Energy Transport

In Lecture 6 we derived three of the four equations of stellar structure. In this lecture we will derive the fourth equation(s), which relates to the transport of energy through the star. The fourth is not one but several equations depending on the source of energy transport (radiation, conduction, or convection).

We will begin by considering the transport of energy via **radiative diffusion**, i.e., heat transfer by random motions. Photons have a mean free path $l = \frac{1}{\sigma n}$ where σ is the cross section and n is the number density of absorbers (or scatterers). Consider the heat flux across a plane at position x with temperature $T(x)$, where the layer at $x-l$ is hotter than the layer at x and the layer at $x+l$ is cooler than the layer at x (see diagram at right). Let u denote the internal energy per unit volume (the energy density).



The heat flux downward is $F_{\text{down}} = \frac{1}{6}vu(x+l)$ where the $1/6$ factor is due to the fact that only one out of six photons are moving down in the x -direction. Notice that the units of F are energy per time per area. The heat flux upward is $F_{\text{up}} = \frac{1}{6}vu(x-l)$. So the net flux in the positive x -direction is: $F_x = F_{\text{up}} - F_{\text{down}} = \frac{1}{6}vu(x-l) - \frac{1}{6}vu(x+l)$. For a mean free path that is small compared to the energy gradient we can approximate: $u(x+l) = u(x) + l du/dx$ and $u(x-l) = u(x) - l du/dx$ which leads us to **Fick's Law**:

$$F_x = -\frac{1}{3}vl \frac{du}{dx} \quad (7.1)$$

This is very general and applies to any case in which the heat transfer occurs via random motions. It applies not only to radiation but also the transport of energy via collisions of particles, also known as **conduction**. The velocity of radiation is c while that of the particles is v ; generally the latter is much smaller than the former, so conduction is almost never important for energy transfer. An important exception is degenerate material, i.e., white dwarfs and degenerate helium cores during post-main sequence evolution, where conduction is a very efficient form of energy transport.

In the specific case of radiative diffusion we can make the following substitutions: $v = c$, $l = \frac{1}{\kappa\rho}$, and $u = aT^4$ and arrive at a fourth equation of stellar structure:

$$\#4a : \quad \boxed{F_{\text{rad}} = -\frac{4acT^3}{3\kappa\rho} \frac{dT}{dr}} \quad (7.2)$$

This equation is used in those regions of a star in which energy transport occurs via radiative diffusion. This equation tells us the temperature gradient necessary to drive a radiative flux F_{rad} through stellar material.

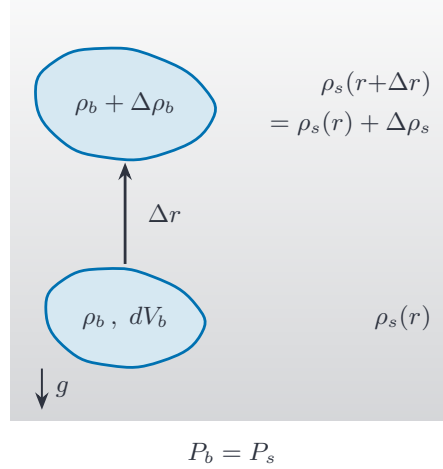
The above is an approximate equation that does not take into account the wavelength-dependence of the radiation field. The more general monochromatic equation is:

$$F_{\text{rad},\nu} = -\frac{4\pi}{3\kappa_\nu\rho} \frac{dB_\nu}{dr} = -\frac{4\pi}{3\kappa_\nu\rho} \frac{\partial B_\nu}{\partial T} \frac{dT}{dr} \quad (7.3)$$

Notice the factor $\frac{1}{\kappa_\nu} \frac{\partial B_\nu}{\partial T}$ – this is the origin of the definition of the Rosseland mean opacity, κ_R , as a weighted (by the Planck function) harmonic mean (recall Lecture 3). The key point is that the frequency-dependence of the heat flux depends inversely on the opacity and on the temperature gradient of the Planck function, and by taking the appropriate integral over frequency one can safely use the earlier (non-monochromatic) version of the radiative diffusion equation by setting $\kappa = \kappa_R$.

The second major source of heat transport in stars is by **convection**. Colloquially, if the heat flux in the star is sufficiently large, then radiative diffusion cannot transport energy quickly enough, and convection will take over. This occurs for example when a pot of water is placed on the stove. Eventually the heat flux is large enough so that the water boils (convection). We will begin by defining a criterion for convective stability, and then go on to describe when a region within a star will be dominated by convective or radiative energy transport, and then we will present simple model for convective energy transport.

We will consider the buoyancy force on a bubble (b) due to the surrounding (s) pressure gradient (see diagram below). The bubble is characterized by a volume, dV_b , and a density,



ρ_b . Importantly, the bubble and its surrounding region have the same pressure: $P = P_b = P_s$ because pressure equalization is almost instantaneous (the timescale for equalization is the bubble size over the sound speed). The (upward, i.e., against gravity) buoyancy force on the bubble is: $F_b = \frac{dP}{dr}dV_b$, where dV_b is the volume of the bubble. In hydrostatic equilibrium the pressure gradient is balanced by gravity: $dP/dr = -\rho_s g$, so $F_b = \rho_s g dV_b$. The net force on the bubble is $F = F_b + F_g = g(\rho_s - \rho_b)dV_b$, or, considering force per unit volume: $f \equiv F/dV_b = g(\rho_s - \rho_b)$.

Consider a bubble initially with $f = 0$, i.e., $\rho_b = \rho_s$. Some perturbation causes a bubble displacement of Δr upward. The new force is $f = g(\Delta\rho_s - \Delta\rho_b)$, where $\Delta\rho$ is the change in the density relative to the initial location. If f has the same sign as Δr , then the blob will feel a net positive buoyancy force and will move in the direction of Δr . This is an instability. The result is convective motions. To make additional progress we need to compute $\Delta\rho_s$ and $\Delta\rho_b$.

In the surroundings, the change in density is just given by the density gradient, i.e., $\Delta\rho_s = \frac{d\rho}{dr}\Delta r$. The change in the bubble density is in principle much more complicated as it depends on the thermodynamics of the bubble. For simplicity, let's assume adiabatic behavior: $d\ln P = \gamma d\ln \rho$, or $\Delta P/P = \gamma \Delta\rho/\rho$ where γ is the adiabatic index. This allows us to write:

$$\Delta\rho_b = \frac{1}{\gamma} \frac{\rho}{P} \Delta P = \frac{1}{\gamma} \frac{\rho}{P} \frac{dP}{dr} \Delta r \quad (7.4)$$

Combining:

$$f = \rho \frac{d^2 \Delta r}{dt^2} = g \left[\frac{d\rho}{dr} - \frac{1}{\gamma} \frac{\rho}{P} \frac{dP}{dr} \right] \Delta r \quad (7.5)$$

or:

$$\frac{d^2 \Delta r}{dt^2} = -N^2 \Delta r \quad , \quad N^2 \equiv g \left[\frac{1}{\gamma} \frac{d\ln P}{dr} - \frac{d\ln \rho}{dr} \right] \quad (7.6)$$

N is known as the Brunt-Väisälä frequency. When $N^2 > 0$ the displacement perturbations are stable and have solutions of the form $\Delta r \propto \sin Nt$, i.e., they are oscillations. These are called gravity waves (not gravitational waves!). When $N^2 < 0$ the perturbations are unstable and solutions have the form: $\Delta r \propto e^{|N|t}$, i.e., the displacements grow exponentially in time.

The case for stability is known as the **Schwarzschild criterion** for convective stability, and can be written as:

$$\frac{1}{\gamma} \frac{d \ln P}{dr} > \frac{d \ln \rho}{dr} \quad (7.7)$$

This criterion can be re-cast in many forms. For example, assuming an ideal gas (and ignoring composition gradients - not a good assumption in many cases!), we can use: $d \ln P / dr = d \ln \rho / dr + d \ln T / dr$ to eliminate ρ and arrive at:

$$\frac{d \ln T}{dr} > \left(\frac{\gamma - 1}{\gamma} \right) \frac{d \ln P}{dr} \quad (7.8)$$

or

$$\frac{d \ln T}{d \ln P} < \frac{\gamma - 1}{\gamma} \quad (7.9)$$

Note that $d \ln P / dr$ is negative, so there was a sign flip in the inequality. Piling on the notation, you will frequently see $\nabla \equiv d \ln T / d \ln P$ and $\nabla_{\text{ad}} \equiv (\gamma - 1) / \gamma$, so the Schwarzschild stability criterion boils down to $\boxed{\nabla < \nabla_{\text{ad}}}$. In words, if the actual temperature gradient in the star is shallower than the adiabatic gradient, then the region is stable to convection.

So when does convection occur? Consider again the energy transport by radiative diffusion: $F_{\text{rad}} = -\frac{4acT^3}{3\kappa\rho} \frac{dT}{dr}$, so $\nabla = -\frac{3\kappa\rho F_{\text{rad}}}{4acT^4} \left(\frac{d \ln P}{dr} \right)^{-1}$ but $\left(\frac{d \ln P}{dr} \right)^{-1} \equiv -H_P$, the pressure scale height. So $\nabla = \frac{3\kappa\rho F_{\text{rad}}}{4acT^4} H_P$. Now let's define:

$$\nabla_{\text{rad}} \equiv \frac{3\kappa\rho F}{4acT^4} H_P \quad (7.10)$$

where F is the total flux. This is the value ∇ would be if all of the energy was transported by radiation. If $\nabla_{\text{rad}} > \nabla_{\text{ad}}$ then the region is **superadiabatic** and convection will occur.

An important feature of ∇_{rad} is that it is sensitive to both the radiative flux and the opacity. A region can therefore be convective if either of these quantities is large. We will find examples of both cases in stellar interiors.

Now we need to actually compute the convective flux. Recall that we have three gradients to consider: ∇ , which is the *actual* gradient in the star, ∇_{ad} , which is the adiabatic gradient, and is solely a function of the thermodynamics of the region, and ∇_{rad} , the gradient *if* the total flux was transported by radiation. So if $\nabla_{\text{rad}} < \nabla_{\text{ad}}$ then $\nabla = \nabla_{\text{rad}}$. Otherwise, we have convection, and the flux is transported by both radiation and convection: $F = F_{\text{rad}} + F_{\text{conv}}$, or $F_{\text{rad}} = F(1 - F_{\text{conv}}/F)$, so $\nabla = \nabla_{\text{rad}}(1 - F_{\text{conv}}/F)$.

The problem we are now faced with is that we don't know how to compute F_{conv} . There is no complete, self-consistent analytic theory, and so 3D (magneto-)hydrodynamical simulations are required. This is a very difficult problem that is an active area of research.

The standard approach to this problem is the **mixing-length theory** (MLT) of convection. The key idea is to assume that a bubble rises a distance $\Delta r = l = \alpha_{\text{MLT}} H_P$ before dissolving into its surroundings. In this formulation, l is the mixing length, α_{MLT} is a dimensionless mixing length parameter of order unity, and H_P is the pressure scale height. The physical picture is of a bubble that rises and expands, decreases in density, and eventually dissolves into the background. The natural length scale in the problem is the pressure scale height, hence the parameterization above.

Our goal is to compute the convective flux, or the energy transported by convection. We can write this as $F_{\text{conv}} = \delta q v_{\text{conv}}$ where δq is the net heat released to the surroundings per unit volume, and v_{conv} is the velocity of the bubble. We can compute δq via $\delta q = C_P \rho (\Delta T_b - \Delta T_s)$, where C_P is the specific heat at constant pressure. Now for some algebra:

$$\frac{\Delta T_s}{T} = \frac{d \ln T}{d \ln r} \frac{\Delta r}{r} = \frac{d \ln T}{d \ln P} \frac{d \ln P}{d \ln r} \frac{\Delta r}{r} = \nabla \frac{r}{H_P} \frac{\alpha_{\text{MLT}} H_P}{r} = \alpha_{\text{MLT}} \nabla \quad (7.11)$$

while for adiabatic bubble motion we have:

$$\frac{\Delta T_b}{T} = \left. \frac{d \ln T}{d \ln P} \right|_{\text{ad}} \frac{d \ln P}{d \ln r} \frac{\Delta r}{r} = \alpha_{\text{MLT}} \nabla_{\text{ad}} \quad (7.12)$$

which leads us to: $\delta q = C_P T \rho \alpha_{\text{MLT}} (\nabla - \nabla_{\text{ad}})$.

We can estimate v_{conv} from the equation of motion for the bubble: $\ddot{r} = -N^2 dr$ and assuming that the final velocity is just $v^2 = 2\ddot{r}l$ after traveling a distance l . Substituting in our definition for N above leads us to

$$v_{\text{conv}}^2 = 2g(\nabla - \nabla_{\text{ad}}) \alpha_{\text{MLT}}^2 H_P. \quad (7.13)$$

Finally:

$$\boxed{F_{\text{conv}} = C_P T \rho \alpha_{\text{MLT}}^2 (\nabla - \nabla_{\text{ad}})^{3/2} (2gH_P)^{1/2}} \quad (7.14)$$

Recall from above $\nabla = \nabla_{\text{rad}}(1 - F_{\text{conv}}/F)$, so we can now solve for both F_{conv} and ∇ .

The superadiabaticity, defined as $\nabla - \nabla_{\text{ad}}$ is generally very small in convective regions, e.g., 10^{-7} , so convection is very efficient at maintaining an adiabatic gradient. There must be *some* non-zero superadiabaticity, otherwise the convective flux would be zero.

Some comments are in order. In general the bubble motion is not adiabatic, and one must account for radiative losses. Close to the stellar surface $v_{\text{conv}} > c_s$, i.e., the bubble motions are supersonic. This signals a breakdown in MLT.

The composition gradient, which we ignored in the derivations above, can in some cases result in lots of additional complexity. For example, a region could be stable according to the Schwarzschild criterion but inclusion of $d\ln\mu/d\ln P$ would imply instability. This case is known as thermohaline mixing. The opposite case, in which the Schwarzschild criterion would suggest instability but composition gradients act to stabilize the region is known as semiconvection. The criterion for stability in the presence of composition gradients is known as the Ledoux criterion.

The exact boundary between radiative and convective zones is uncertain and not predicted within the mixing length theory. In this context one discusses “convective overshoot” in which convective bubbles overshoot into the formally stable regions. Recall that our derivation of the condition for stability referred to forces, not velocities. The issue of overshoot is very important both from the point of view of stellar mixing and the availability of fresh fuel for nuclear burning in the cores of massive stars, as we will see later. The main energy transport is via radiation in the overshooting region.

In spite of these many shortcomings, nearly all modern 1D stellar evolution codes (including MESA) adopt MLT for the physical description of convective regions in stars. The single free parameter, α_{MLT} , is usually tuned to reproduce key properties of the Sun, and is then fixed for all other stars. Recent 3D simulations suggest that α_{MLT} may vary by modest amounts as a function of temperature, metallicity, and surface gravity (e.g., Magic et al. 2014).

In addition to being important for the transfer of energy, convection is an efficient mechanism for **mixing**. In general, a region that is experiencing convection will have a flat composition gradient. The implications of this for stellar evolution will appear in later lectures.

Summary: We have now completed the assembly of the key equations of stellar structure. By far the most significant source of uncertainty in these equations is the treatment of convection, which is intrinsically a 3D process that we have attempted to describe in 1D. It should not surprise you to learn that the treatment of convection is one of the most important outstanding uncertainties in the theory of stellar structure and evolution. Modern 3D stellar evolution codes can only evolve a star for a small fraction of its lifetime, and so one must take insight from the 3D models and deploy them in 1D programs such as MESA.

Chapter 8

Simple Stellar Models

We have now completed assembly of the key equations of stellar structure and described the most important microphysics related to stars. Much of the rest of this course is devoted to understanding the implications of these equations. In fact, it is somewhat remarkable that these straightforward equations give rise to such a diverse range of stellar types and evolutionary pathways.

In this lecture we will provide an overview of how one solves these equations and then define and explore the behavior of the stellar main sequence. We will also sketch a derivation of the Eddington limit.

By way of summarizing the past several lectures, the fundamental equations of stellar structure are:

$$\frac{dP}{dr} = -g\rho \quad , \quad g = \frac{Gm}{r^2} \quad (8.1)$$

$$\frac{dm}{dr} = 4\pi r^2 \rho \quad (8.2)$$

$$\frac{dL}{dr} = 4\pi r^2 \rho \epsilon \quad (8.3)$$

$$L = 4\pi r^2 (F_{\text{rad}} + F_{\text{conv}}) \quad (8.4)$$

where: $F_{\text{rad}} = -\frac{4acT^3}{3\kappa\rho} \frac{dT}{dr}$ and $F_{\text{conv}} = C_P T \rho \alpha_{\text{MLT}}^2 (\nabla - \nabla_{\text{ad}})^{3/2} (2gH_P)^{1/2}$ if $\nabla > \nabla_{\text{ad}}$, otherwise $F_{\text{conv}} = 0$. In the above equations we have assumed hydrostatic equilibrium.

In order to solve these equations we require a few more pieces of information. First, the “constitutive relations” (equations that pertain to materials or substances), which include the equation of state (EOS), the opacity, and the energy sources. In general, $P =$

$P(\rho, T, \text{composition})$, $\kappa = \kappa(\rho, T, \text{composition})$ and $\epsilon = \epsilon(\rho, T, \text{composition})$. These relations are very complicated and modern programs use large tabulated values for these. In this course we will make the assumption in most cases that the EOS is a combination of an ideal gas law and radiation pressure:

$$P = P_{\text{gas}} + P_r = \frac{\rho k T}{\mu m_{\text{H}}} + \frac{a T^4}{3} \quad (8.5)$$

Finally, we require **boundary conditions**. Several are trivial, e.g., $m(0) = 0$, $L(0) = 0$, and $m(R) = M$. Others are the subject of much work, in particular the outer boundary conditions for T and P . As a first attempt we might set $P(R) = 0$ and $T(R) = 0$, but this is not quite right, as the “surface” of the star has $T(R) = T_{\text{eff}}$, and $P(R) = P_{\text{phot}}$ where P_{phot} is the pressure at the photosphere. State of the art models use tabulated atmosphere boundary conditions based on sophisticated radiative transfer calculations.

These equations are deceptively simple; much is hidden behind the chemical composition. For example, changes in composition are the ultimate driver of stellar evolution. We need to therefore track and evolve the abundances. As an example, on the main sequence we have: $dX/dt = -\epsilon_{\text{nuc}}/26.7 \text{ MeV}$, $dY/dt = -dX/dt$, and $dZ/dt = 0$ (approximately). In other words, Hydrogen is burned into helium at a rate given by ϵ . Additional equations are required to capture certain processes, such as element diffusion and angular momentum transport.

It is sometimes said that stellar evolution depends on only two parameters, mass and composition. This is referred to as the Vogt-Russell theorem. While this is a good rule of thumb, it is not correct in all cases, even with the simplified assumptions we have made here.

We have neglected multiple important physical effects including mass-loss, rotation, stellar binarity, and magnetic fields. We will discuss these topics later in the course.

There exists a large literature devoted to numerical solutions to these equations. The basic approach is to treat the problem in 1D and to place the star on a grid (in either radius or mass). One can then turn the differential equations into finite differences, e.g. dP/dr becomes $\frac{P_{i+1} - P_i}{r_{i+1} - r_i}$. The problem can become quite complex, as modern models such as **MESA** allow for variable timesteps and on-the-fly adjustment of the sizes and number of mesh points. For those familiar with hydro codes, **MESA** is basically a 1D AMR (adaptive mesh refinement) code.

In principle, we’re done! Much of the rest of this course consists of “applications” of these equations, and we could simply turn to **MESA** to solve the equations. However, this would not provide us with an *understanding* of the resulting behavior. A guiding principle going

forward will be to extract intuition for the general behavior without relying on the full solution. Sometimes this is straightforward (and enlightening), while in other cases an intuitive explanation of the full solutions is surprisingly elusive.

Lets begin at the **main sequence**. The so-called zero-age main sequence (ZAMS) represents the time when stars are still chemically homogenous, i.e., just before H fusion becomes the dominant energy source. It is important to have a strong intuitive grasp of the main sequence because stars spend the majority of their lives ($\sim 90\%$) in this phase. It is called a main sequence because stars in this phase occupy a narrow sequence in the HR diagram. The sequence is controlled by the stellar mass.

We will frequently refer to a core-envelope structure within stars; we will see examples of how and why this structure emerges throughout the course. Examples include stars with convective cores and radiative envelopes and stars with radiative cores and convective envelopes.

There are several important transition points along the main sequence (keep in mind that the exact boundaries are approximate and depend on multiple factors including composition):

- $\approx 0.1M_{\odot}$: below this mass an object will not fuse H and so is not a star (objects below this mass are known as brown dwarfs). See Lecture 9 for details.
- $\approx 0.3M_{\odot}$: below this mass the star is fully convective, while above this mass the core is radiative while the envelope is convective.
- $\approx 1.5M_{\odot}$: above this mass the CNO cycle is the dominant source of H fusion. The very steep temperature dependence results in a convective core. Below this mass the pp chain is dominant, and produces a radiative core.
- $\approx 50M_{\odot}$: radiation pressure becomes a major source of pressure above this mass.

The origin of the $1.5M_{\odot}$ transition point can be understood as follows. For low-mass stars $\epsilon \propto T^4$, so energy generation is somewhat spread out in the inner region of the star. This implies that $F = L/4\pi r^2$ is not very large, and hence $\nabla < \nabla_{\text{ad}}$. In contrast, for higher-mass stars we have $\epsilon \propto T^{17}$, so energy generation is very concentrated, F is large, $\nabla > \nabla_{\text{ad}}$, and so convection develops in the core. Because the core is convective, the composition gradient is flat. The core will therefore evolve in a different way for low and high mass stars (as we will discuss in Lecture 10).

As the star evolves, μ increases in the core. At the same time, the central temperature, T_c , is nearly constant during H burning (nuclear burning acts like a thermostat because of the high temperature sensitivity of the reaction rates). The ideal gas law applied to the core implies $P_c/\rho_c \sim T_c/\mu_c$. Since T_c is constant and μ_c is increasing, this means that the density must increase faster than the pressure to maintain hydrostatic equilibrium. An increase in density can come about only if the core shrinks, which it will do during the main sequence.

We are now going to make some back of the envelope estimates of how various quantities scale with one another along the main sequence. We will take the approach of dropping constants and replacing dx/dy with x/y (for arbitrary variables x and y). This is obviously a major approximation, but it turns out to work in obtaining the correct scalings. To indicate that we're making very rough approximations (including dropping factors of 4π , etc.), we're going to use " $\sim$ " below.

Lets consider models in which radiation is the dominant energy transport mechanism. The fourth equation of stellar structure becomes:

$$F_{\text{rad}} \sim \frac{L}{R^2} \sim \frac{acT^3}{\kappa\rho} \frac{T}{R} \quad (8.6)$$

rearranging:

$$L \sim (aT^4) (R^3) \left(\frac{Rc}{\kappa M} \right) \quad (8.7)$$

this is an energy density times a volume divided by a diffusion timescale, i.e., an energy divided by a timescale, i.e., a luminosity. To see why the last term can be viewed as a diffusion timescale, consider that the diffusion timescale is in general $t_{\text{diff}} \sim R^2/D$, where D is the diffusion constant and R is a length scale. In general, the diffusion constant scales as $D \sim vl$ where l is the mean free path. For photons diffusing in a star, $v = c$ and $l = 1/\kappa\rho \sim R^3/\kappa M$. Combining these pieces leads to: $t_{\text{diff}} \sim \frac{\kappa M}{Rc}$.

Question: Notice a remarkable feature of Equation 8.7 - it does not depend on the energy production mechanism! How can this be?

We can make further progress by using the ideal gas law to eliminate T : $P_{\text{gas}} \sim \frac{kMT}{\mu m_{\text{H}} R^3} \sim GM^2/R^4$ so $TR \sim \frac{GM\mu m_{\text{H}}}{k}$. What temperature does T refer to? Its hard to say since this is all very qualitative, but it is closer to the central temperature, T_c than the surface (it is the central pressure that is providing the support against gravity). Using this equation to write

the luminosity without reference to temperature leads us to:

$$L \sim \frac{ac}{\kappa} \left[\frac{G\mu m_{\text{H}}}{k} \right]^4 M^3 \quad (8.8)$$

We now can see why stars span such a wide range in luminosity ($\sim 10^9$) despite only spanning a factor of 10^3 in mass.

For stars with $M > 2M_{\odot}$, where $\kappa \sim \text{constant}$ due to electron scattering (but not too high mass so that gas pressure is still the dominant source of pressure), we have: $L \propto \mu^4 M^3$.

For lower-mass stars ($\lesssim 2M_{\odot}$) the opacity is governed by Kramer's Law: $\kappa \propto \rho T^{-7/2}$. We also need to know that the radius depends weakly on mass in this regime: $R \sim M^{1/2}$ (see LL Ch 13.1.2 for details). Combining: $L \propto \mu^{7.5} M^{5.25}$.

At the highest masses ($M \gtrsim 50M_{\odot}$) radiation pressure dominates and the opacity is provided by electron scattering, so $P \sim T^4 \sim M^2/R^4$ and hence $T^4 R^4 \sim M^2$. Returning to our earlier equation for the luminosity, which was: $L \sim aT^4 R^3 Rc/\kappa M \sim T^4 R^4/M$, leads us to $L \sim M$. This is the regime of radiative-pressure supported stars with energy transport governed by radiative diffusion.

At the lowest masses stars are fully convective and the relations derived here, which assumed energy transfer via radiation, will not apply. Such stars will be discussed in Lecture 9.

Question: What do these relations between L and M imply about the main sequence lifetime as a function of mass?

We're now going to derive the Eddington limit, a very important concept that sets a theoretical upper limit to the luminosity of an object.

From the fourth equation of stellar structure for radiation we have:

$$F_{\text{rad}} = -\frac{c}{\kappa\rho} \frac{dP_{\text{rad}}}{dr} \quad (8.9)$$

We can recast the equation for hydrostatic equilibrium as:

$$\frac{dP}{dr} = \frac{dP_{\text{gas}}}{dr} + \frac{dP_r}{dr} = -g\rho \quad (8.10)$$

combining:

$$\frac{dP_{\text{gas}}}{dr} - \frac{\kappa\rho F_{\text{rad}}}{c} = -g\rho \quad (8.11)$$

rearranging:

$$\frac{dP_{\text{gas}}}{dr} = -g\rho \left[1 - \frac{\kappa F_{\text{rad}}}{gc} \right] \quad (8.12)$$

The term in brackets must be positive, and so we can define $\Lambda \equiv \frac{\kappa F_{\text{rad}}}{gc}$ and require $\Lambda < 1$. Λ is basically the ratio of the radiation pressure gradient to the gravitational restoring force. At the surface (where $g = GM/R^2$ and $F = L/4\pi R^2$), we require:

$$\Lambda = \frac{\kappa L}{4\pi GMc} < 1 \quad (8.13)$$

Setting $\Lambda = 1$ leads to the definition of the **Eddington Limit**:

$$L_{\text{Edd}} \equiv \frac{4\pi GMc}{\kappa} \quad (8.14)$$

A star (or any object) cannot be in hydrostatic equilibrium if it has a luminosity in excess of L_{Edd} . At high temperature the opacity is governed by Thomson scattering, so $\kappa = \sigma_T/m_H$ and the equation for the Eddington luminosity is particularly simple: $L_{\text{Edd}} = 3 \times 10^4 L_{\odot} (M/M_{\odot})$.

If $\Lambda > 1$ then radiation pressure overwhelms gravity. What happens? In 1D models the program usually crashes. In the real world, the star has to adjust/react somehow. We don't really know what happens, but various instabilities likely develop, and mass is probably lost from the star. At least in models, stars can have $\Lambda \approx 1$, not only at the surface but deeper in the interior. The high values of Λ are often driven by sudden increases in the opacity, due to particular ionization states of atoms or molecules becoming important at specific temperatures.

The Eddington limit also sets an upper bound on how rapidly black holes can grow, as we will see in Lecture 19.

Summary: In this lecture we have sketched out the interior structure of stars as well as their behavior along the main sequence. It is important to understand the different regimes along the main sequence, including the dominant pressure source and energy transport and production mechanisms, and the nature of the core. These shape not only the observables and lifetimes of stars on the main sequence but also play a crucial role in the subsequent post-main sequence evolution, as we will see in later lectures.

Chapter 9

Fully Convective Stars

In this lecture we will consider the behavior of fully convective objects. This includes the important concept of the Hayashi limit, and the evolution of protostars through the pre-main sequence phase. We will also consider the physics setting the minimum stellar mass.

There are several examples of fully (or near fully) convective stars in nature, including stars in their pre-main sequence contraction phase, very low mass stars, and evolved low-to-intermediate mass stars. Unlike the previous lecture in which we appealed to energy transport by radiation, the luminosity of a fully convective star is not determined by its interior structure. It must instead be set by the surface conditions, in particular the opacity at the surface. A key concept in this context is the Hayashi limit, which defines a nearly vertical track in the HR diagram. Stars cannot exist (in hydrostatic equilibrium) cooler than the Hayashi limit. Let's see why.

Recall from Lecture 3 that a simple plane-parallel stellar atmosphere is governed by: $dP/d\tau = g/\kappa$, where τ is the optical depth. Let's assume κ is independent of τ (depth), and recall g is (to a good approximation) constant throughout the atmosphere. If we define the “surface” at $\tau = 2/3$ (see HW1, problem 3), then we have $P_{\text{surf}} = \frac{2GM}{3\kappa R^2}$. Let's assume the surface is relatively cool $T_{\text{eff}} < 4000$ K, where ideal gas pressure dominates and the opacity is dominated by the H^- molecule. We can then set $P_{\text{surf}} = \frac{k}{\mu m_H} \rho_{\text{surf}} T_{\text{eff}}$ and the H^- opacity scales approximately as: $\kappa = \kappa_0 \rho^{1/2} T^{7.7}$. The very steep temperature dependence of κ will be a key part of this story. Combining these two expressions for the surface pressure leads to:

$$\rho_{\text{surf}}^{1.5} T_{\text{eff}}^{8.7} = \frac{2}{3} \frac{\mu m_H}{k} \frac{GM}{\kappa_0 R^2} \quad (9.1)$$

The next step is to appeal to the $n = 1.5$ ($\gamma = 5/3$) polytropic solution appropriate for a fully

convective star (see e.g., Table 4.1 in Pols). This provides us with the following relations: $\rho/T^{1.5}=\text{constant}$ (and hence $\rho_{\text{surf}}/T_{\text{eff}}^{1.5} = \rho_c/T_c^{1.5}$), $\rho_c = 6\bar{\rho} = 6\frac{3M}{4\pi R^2}$, and $T_c = 0.54 \frac{\mu m_H}{k} \frac{GM}{R}$. Finally, we will use $L = 4\pi R^2 \sigma T_{\text{eff}}^4$ to remove R from the resulting expression:

$$T_{\text{eff}}^{11.5} \approx \frac{0.07}{\kappa_0 \sigma^{1/8}} \left(\frac{\mu m_H G}{k} \right)^{3.25} M^{1.75} L^{1/8} \quad (9.2)$$

or

$$T_{\text{eff}} \approx 2000\text{K} \left(\frac{M}{M_\odot} \right)^{0.15} \left(\frac{L}{L_\odot} \right)^{0.01} \left(\frac{Z}{0.02} \right)^{-0.04} \quad (9.3)$$

where here we have inserted constants and the fact that $\kappa_0 \propto Z^{1/2}$. This equation describes the **Hayashi limit**, and represents a nearly vertical line in the HR diagram. There should be no stars to the right of the Hayashi limit as they would not be in hydrostatic equilibrium. The limit moves to cooler temperatures for more metal-rich stars.

Intuitively, we can understand the steepness of the Hayashi line in the following way. Suppose a fully convective star increases its radius slightly while attempting to keep L constant. Then the temperature of the photosphere would decrease and the photosphere would become much more transparent (because the H^- opacity is such a strong function of temperature). Hence energy can more easily escape from the interior, which in turn means that the luminosity will greatly increase, even with a slight decrease in photospheric temperature.

How could a star find itself on the cool side of the Hayashi line, and what would happen in that case? The Hayashi line is a locus of $\nabla = \nabla_{\text{ad}}$ (recall that convection is very nearly adiabatic). It turns out that stars can lie cooler than the Hayashi line if $\nabla > \nabla_{\text{ad}}$ (this is not obvious and requires actual calculations). In this case, the convective flux would be extremely large and would result in extremely large luminosities. Such a large energy flux will rapidly (on a dynamical timescale) transport the heat outwards until $\nabla = \nabla_{\text{ad}}$ is re-established. The star will therefore quickly readjust to lie on the Hayashi line.

In our earlier discussion of polytropes we noted that fully convective stars are $n = 3/2$ polytropes (recall $P = K\rho^{1+1/n}$), just like non-relativistic degeneracy pressure. In the latter case it was stated that such objects obey $K = C M^{1/3} R$ with K a constant set by fundamental variables (C is another constant). Such objects have the unusual feature that their radius decreases as their mass increases. Does this also occur for fully convective stars? The answer is no, because the variable K is not a fundamental constant. Instead, for such stars K will be different for different stars, even while $n = 3/2$. K is related to the entropy of the star.

We're now going to discuss the road to the main sequence, including a (very) cursory overview of star formation.

The collapse of a gas cloud occurs when $|\frac{1}{2}U| > K$ where U and K are the potential and kinetic energy (total, not per unit volume). For a spherical cloud with uniform density we have

$$U = -\frac{3}{5} \frac{GM^2}{R} \quad , \quad K = \frac{3}{2} kT \frac{M}{\mu m_{\text{H}}} \quad (9.4)$$

So collapse will occur if

$$M > \frac{5kTR}{\mu m_{\text{H}}G} \equiv M_J \quad (9.5)$$

where M_J is known as the **Jeans Mass**. We can cast this another way by noting that $R \propto (M/\rho)^{1/3}$ so $M_J \propto T^{3/2} \rho^{-1/2}$. We also know that $kT \propto \sigma^2 \sim c_s^2$ where σ is the mean squared speed and c_s is the sound speed. So $M_J \sim c_s^3 \rho^{-1/2}$. A simpler way to understand the Jeans mass can be obtained by considering the length scale over which a sound wave can travel in a dynamical time. The Jeans mass is the mass enclosed within that length scale.

Cooling is efficient during the early stages of cloud collapse because the gas is still optically thin. This implies that the collapse will proceed at approximately constant temperature, i.e., $kT \sim c_s^2 \sim \text{constant}$. As the cloud collapses ρ increases, so M_J decreases. This leads to **fragmentation** of the initial cloud into many smaller collapsing clouds. In other words, a single large cloud will give rise to many small clouds, i.e., many proto-stars. Eventually the density is so high that cooling is no longer efficient and collapse continues adiabatically. In this regime the temperature rises as the collapse continues, which cause the Jeans mass to *increase*. Fragmentation stops, and we are left with a number of clumps within the original cloud.

These clumps continue to collapse (if above a minimum mass threshold, see LL Ch 12.5 for details), and clump collapse ends when the clump temperature is sufficiently high so that the pressure counteracts gravity and hydrostatic equilibrium is achieved. This occurs when molecular hydrogen is dissociated and H is subsequently ionized – when this occurs cooling within the clump becomes inefficient and the temperature rises substantially. Lamers and Levesque work out the details (in Ch 12.6), and the result is a relation between mass and radius of the form:

$$R/R_{\odot} \approx 100M/M_{\odot} \quad (9.6)$$

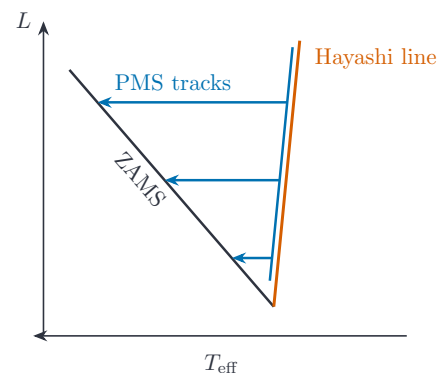
These are **protostars**, and they are fully convective (they have high opacities), hence have $T \sim 3000$ K (via Hayashi), hence $L \sim 10^3 L_{\odot}$ for $M = 1M_{\odot}$ (assuming $L = 4\pi R^2 \sigma T^4$).

Protostars are very luminous! The luminosity is sourced by gravitational contraction. At this point the mean temperature is $3/2 M/\mu m_{\text{H}} k \bar{T} = 5/2 GM^2/R$, i.e., $\bar{T} \sim 7 \times 10^4$ K. Too low for fusion. The **pre-main sequence** phase begins when the luminosity is due to gravitational contraction and when the evolutionary timescale is governed by the Kelvin-Helmholtz timescale.

The collapse is reducing the radius by several orders of magnitude, so a tiny initial rotation will be greatly amplified by conservation of angular momentum. The final configuration will almost certainly involve a protostellar disk. This is ignored here, but is essential for the formation of planets, etc.

The protostar is radiating away energy, so the star must contract to maintain hydrostatic equilibrium. The star will continue to contract until the mean temperature exceeds $\sim 10^6$ K, at which point the opacity has dropped dramatically so that the star becomes radiative, at least in the core. Because $\bar{T} \sim M/R$ (assuming an ideal gas), and we're appealing to a $\sim 50\times$ increase in temperature, this implies a radius decrease of $50\times$, i.e., $2R_{\odot}$ for $1M_{\odot}$. The star is still on the Hayashi line, so $T_{\text{eff}} \sim 3000$ K and $L \sim 0.5L_{\odot}$. The timescale for this to occur is the Kelvin-Helmholtz time: $t_{\text{KH}} = |E_{\text{pot}}|/L$. We need to be careful in how we appeal to KH because R and L are changing a lot. So one can integrate in time or take some "suitable" average. The result of this (or by appeal to more sophisticated models) is a timescale of 10^6 yr for $1M_{\odot}$. The mass scaling can be estimated as follows. $t \sim M^2/LR \sim M^2/R^3$ (assuming T_{eff} constant). We had above that $R \sim M$ (see again LL Eqn 12.13), so $t \sim M^{-1}$. This is the timescale to descend along the Hayashi line. The star leaves the Hayashi line when its internal temperature has risen to the point (and opacity has dropped to the point) that the star is no longer fully convective.

The next phase is evolution from the Hayashi limit to the zero-age main sequence (ZAMS), also known as pre-main sequence contraction (see diagram at right). The protostar departs the Hayashi limit when the core is radiative, so our relations from earlier lectures apply: $L \propto \mu^a M^b$. This means that the evolution of a star from the Hayashi limit to ZAMS occurs at fixed luminosity, i.e., horizontally in the HR diagram. The timescale for this evolution is again the relevant KH time: $t \sim |E|/L \sim 6 \times 10^7 (M/M_{\odot})^{-2.5}$ yr. At $M = 0.3M_{\odot}$ the timescale is $\sim 10^9$ yr – we should therefore expect to observe many low-mass stars in their pre-main sequence phase. The



Comment: The evolution from protostar to ZAMS involves the interplay of multiple physical processes that we have covered so far in class (convection vs. radiation, core vs. envelope, dynamical vs. thermal vs. nuclear timescales, gravitational vs. nuclear energy sources). You will work through MESA calculations of this phase in HW3. Take the time to make sure you understand this phase, both for its own sake and as an example of how to reason through the relative importance of the various physical processes.

timescale for this phase is much longer because the luminosity is much lower, compared to the descent along the Hayashi line. Realistic models predict that the luminosity increases by a factor of $2 - 3\times$ during this phase because the central temperature is increasing, which results in decreasing opacity. In the above scaling relations for L we had implicitly been assuming $\kappa = \text{constant}$, but in fact (see Lecture 8), $L \propto \kappa^{-1}$. So a decreasing opacity results in an increasing luminosity, if everything else is held constant. This phase ends when T_c reaches the temperature where fusion occurs.

Some interesting things happen during pre-main sequence contraction. Big Bang Nucleosynthesis (BBN) produces a small amount of initial cosmic D (^2H) and ^7Li . These elements burn at relatively low temperature. At $\sim 10^6$ K the following reaction occurs:



and at $\sim 3 \times 10^6$ K:



Li in particular is a novel probe of pre-main sequence evolution. Of course, we can only observe elements on the surface, so we will only see depletion of Li due to burning if convection reaches the core, where $T_c \sim 3 \times 10^6$ K. This only occurs for $M \lesssim 1.4M_\odot$. For more massive stars, they are radiative in their outer envelopes by the time the core reaches the threshold for Li burning.

The so-called lithium depletion boundary (LDB) is a unique and sensitive probe of a stellar population's age. Lower-mass stars take longer to reach $T_c \sim 3 \times 10^6$ K, so the luminosity where stars still have Li on the surface is a chronometer (one of the few that is not particularly sensitive to the details of stellar models). The primary limitation is observational: measuring the weak Li I line at $6708\text{\AA}$ in faint low-mass stars.

We're now going to consider the physics determining the **minimum stellar mass** for nuclear burning.

Recall that the core temperature can be estimated by $T_c \approx \frac{\mu m_H}{k} \frac{GM}{R}$. In our case with a contracting core in which the core supplies most of the pressure, we have: $T_c \approx \frac{\mu m_H}{k} \frac{GM_c}{R_c}$, or replacing radius with density:

$$T_c \approx \frac{\mu m_H}{k} G M_c^{2/3} \rho_c^{1/3} \quad (9.9)$$

The core needs to contract until $T_c \approx T_{\text{ign}}$ for fusion. But if M_c is too low, then the density will reach conditions for electron degeneracy before this can happen. We can estimate when this will occur by setting $P_{\text{ideal}} = P_{\text{degen,e}}$, which leads to the following condition:

$$\frac{k}{\mu_c m_H} \rho_c T_c = K (\rho_c / \mu_e)^{5/3} \quad (9.10)$$

Re-arranging:

$$T_c = \frac{K m_H}{k} \left(\frac{\mu_c}{\rho_c} \right) \left(\frac{\rho_c}{\mu_e} \right)^{5/3} \quad (9.11)$$

The combination of Eqns 9.9 and 9.11 leads to the condition for the core mass, below which degeneracy pressure dominates and the core temperature will cease rising as contraction continues. Specifically, we can use these two relations to eliminate ρ_c and derive a relation between M_c and T_c :

$$M_{c,\text{min}} = \left[\frac{kK}{\mu_c \mu_e^{5/3} m_H G^2} \right]^{3/4} T_{\text{ign}}^{3/4} \quad (9.12)$$

If we set $T_{\text{ign}} = 6 \times 10^6$ K for H ignition, this leads to $M_{\text{min}} \approx 0.1 M_\odot$. Setting $T_{\text{ign}} = 10^9$ K for He ignition (along with $\mu_e = \mu_c = 2$ for a He core) leads to $M_{c,\text{min}} \approx 1.0 M_\odot$. This latter boundary is a bit more approximate. In reality, stars around $1 M_\odot$ do burn He in their centers, but lower-mass stars do not (or rather will not - stars with masses low enough to fail to burn He have not yet left the main sequence).

Notice that the minimum mass depends on composition through the mean molecular weight. Our sketch refers explicitly to the “core” and not the entire star and makes several significant approximations, but it does get close to the correct answer. Modern calculations place the minimum H burning mass at $0.07 - 0.08 M_\odot$ (e.g., Chabrier et al. 2000). Objects with masses below this threshold never become stars. They are instead known as **brown dwarfs**, are supported by electron degeneracy pressure, and remain very cool and dim. As in the case of white dwarfs, brown dwarfs have an inverted mass-radius relation (larger masses result in smaller sizes). A Jupiter mass, M_J , is approximately $0.1\% M_\odot$, so that a brown dwarf with mass $M = 0.05 M_\odot = 50 M_J$. The brown dwarf regime is bracketed by stars on one end and gas giant planets on the other. The boundary between brown dwarfs and giant planets is defined by the former’s ability to burn deuterium. Deuterium burning occurs over the mass range $\approx 13 - 80 M_J$. At $\gtrsim 65 M_J$ brown dwarfs are also able to burn lithium.

Chapter 10

Post-Main Sequence Evolution

In the first nine lectures the narrative was relatively linear and straightforward. Behavior along the main sequence could be isolated into several key regimes, each governed by one or two basic principles (e.g., the balance of radiative vs. convective energy transport). Once a star departs the main sequence, the story becomes dramatically more complicated. The emergent behavior is in many cases quite sensitive not only to the initial conditions (e.g., mass and composition), but also the parameterization of the (1D) model, e.g., via the treatment of convective overshooting or stellar mass loss, and even to uncertain nuclear reaction networks.

Nevertheless, there are some guiding principles that we will rely on as we seek an intuitive understanding of post-main sequence evolution. The goal of lecture is to explore some of those basic principles.

As we remarked earlier, stars frequently develop a core-envelope structure. In some phases such as the main sequence, the core is undergoing nuclear fusion, which provides the pressure necessary to forestall gravitational collapse. When a period of core burning comes to an end we often have an **isothermal core**. This occurs because $\epsilon = 0$, so $F_{\text{rad}} \propto dT/dr = 0$, implying $dT/dr = 0$. The core will undergo a period of contraction (due to the loss of pressure support) until the core-envelope boundary is hot enough for fusion to occur. This is known as **shell burning**. As shell burning proceeds, the ashes contribute to a growing core below. During this process the core will increase in density, in some cases reaching a density great enough to undergo the next phase of core burning. This process of alternating core and shell burning can proceed in several cycles (from H core to H shell, to He core to

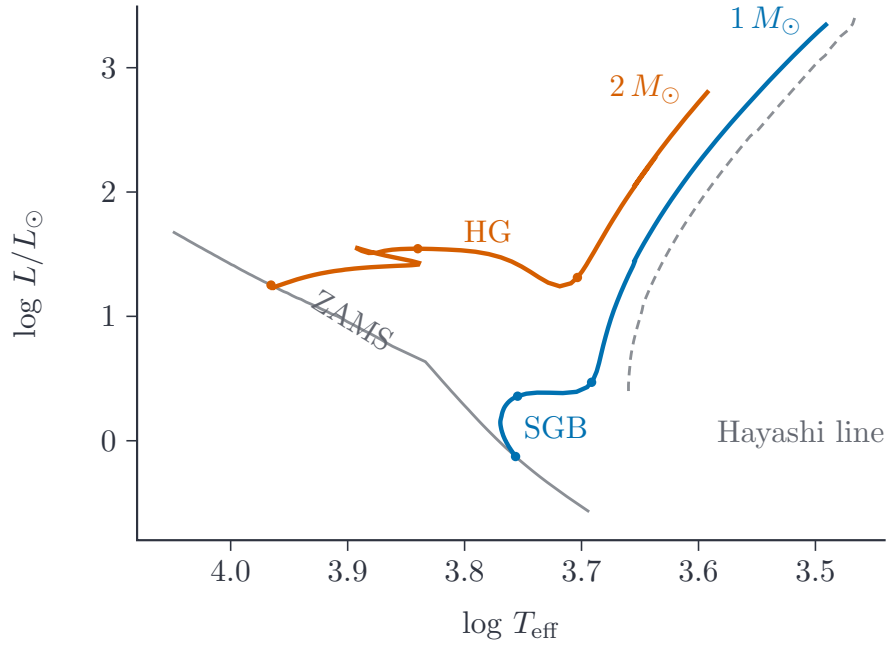
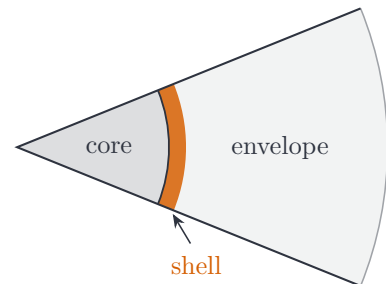


Figure 10.1: Post-main sequence evolution of a $1M_{\odot}$ and a $2M_{\odot}$ star from the ZAMS to the RGB tip; MIST models at solar metallicity (Choi et al. 2016). Dots bracket the subgiant branch (SGB), crossed in 835 Myr, and the Hertzsprung gap (HG), crossed in 19 Myr.

He shell burning), depending on mass. There are many interesting details to this general picture that we will see along the way.

Stars above $\approx 1.5M_{\odot}$ have convective cores, while lower-mass stars have radiative cores. One important feature of a convective core is that it will have a flat composition profile (due to convective mixing). This means that the end of H fusion is abrupt, and results in a fairly rapid contraction of the core until shell burning is initiated. This creates a distinctive “hook” feature in the HR diagram (sometimes referred to as the Henyey hook, see Figure 10.1) in post-main sequence evolution of stars with $\gtrsim 1.5M_{\odot}$. This feature is also observed in star clusters, and serves as an important constraint on the mass at which stars develop convective cores.



One question we need to address is whether or not the core can support the weight (pressure) of the overlaying envelope. To explore this, let’s consider an isothermal core with temperature T_c . The pressure from the envelope at the core-envelope boundary is P_{env} . We can apply the **virial theorem** to the core: $4\pi R_c^3 P_c - 2K_c = U_c$, where P_c is the pressure at the outer

boundary of the core. This form of the VT might look different than you've seen before. The first term is the external pressure, which in many astronomical applications is zero. If we assume an ideal gas, then $K_c = \frac{3}{2} \frac{M_c}{\mu_c m_H} kT_c$, and $U_c = -\frac{3}{5} \frac{GM_c^2}{R_c}$. Combining:

$$4\pi R_c^3 P_c - 3 \frac{M_c}{\mu_c m_H} kT_c = -\frac{3}{5} \frac{GM_c^2}{R_c} \quad (10.1)$$

solving for P_c :

$$P_c = \frac{3}{4\pi R_c^3} \left[\frac{M_c kT_c}{\mu_c m_H} - \frac{GM_c^2}{5R_c} \right] \quad (10.2)$$

The pressure at the outer boundary of the core, P_c , must equal P_{env} for hydrostatic stability. This is achieved by adjusting R_c as M_c increases. The problem is that $P_c(R_c)$ has a maximum. we can find the maximum by setting $dP_c/dR_c = 0$, which leads to $R_c^{\text{max}} = \frac{4}{15} \frac{GM_c \mu_c m_H}{kT_c}$, or $P_c^{\text{max}} = \frac{10}{G^3 M_c^2} \left(\frac{kT_c}{\mu_c m_H} \right)^4$. So as M_c increases (e.g., due to shell burning growing the core), eventually we will have $P_c^{\text{max}} < P_{\text{env}}$ and the core will collapse. This is known as the **Schönberg–Chandrasekhar Limit** (no relation to the Chandrasekhar mass).

It will be more informative to recast the SC limit in terms of a relation between the core mass, M_c , and total mass, M . To make progress we need to know the pressure at the base of the envelope. We know that the pressure can generally be expressed as GM^2/R^4 to within a factor of a few. In our case we might imagine setting R to the pressure scale height (or a multiple thereof). We can appeal to the virial theorem at the base of the envelope (ignoring factors of order unity): $kT_{\text{env}} \approx kT_c \approx GM\mu_{\text{env}}m_H/R$, where the first equality reflects the fact that the temperature is continuous across the core-envelope boundary. Using the above expression to remove R from our guess of the envelope pressure, we have:

$$P_{\text{env}} \approx \frac{GM^2}{R^4} \approx \left(\frac{kT_c}{\mu_{\text{env}} m_H} \right)^4 \frac{1}{G^3 M^2} \propto \frac{T_c^4}{M^2} \quad (10.3)$$

setting $P_c^{\text{max}} = P_{\text{env}}$ and doing proper calculations for the constants leads to a more useful form of the SC limit:

$$\frac{M_c}{M} \approx 0.37 \left(\frac{\mu_{\text{env}}}{\mu_c} \right)^2 \quad (10.4)$$

For a solar metallicity star leaving the main sequence, $\mu_c = 1.3$ (ionized He) and $\mu_{\text{env}} \approx 0.6$ (ionized solar composition), leading to $M_c \approx 0.08M$. If the core grows beyond this limit then it will collapse (until something else halts the collapse, such as the onset of degeneracy pressure or He fusion).

Some comments regarding the SC limit:

- A contracting core results in the generation of heat and a temperature gradient. This gives rise to a radiating core. This happens on a Kelvin-Helmholtz (KH, or thermal)

timescale, which remember from Lecture 1 is generally short compared to the nuclear timescale. This is the origin of the **Hertzsprung Gap** in the HR diagram, in which very few (no) stars are observed in between the main sequence and the red giant branch (RGB) at $> 1.5M_\odot$.

- At lower masses $\lesssim 1.3M_\odot$ the He core is partly degenerate, and so the SC limit does not apply (we assumed an ideal gas EOS in the derivation of the SC limit). For such stars the evolution from the main sequence to the RGB is slower and there is no gap. This phase of evolution is known as the subgiant branch (SGB), and is well-populated in observations.

Let's consider one specific example of a $5M_\odot$ star evolving off the main sequence (see Pols Section 10.2.1). During H burning the core slowly contracts and the radius increases, resulting in lower T_{eff} . When H is exhausted in the core, the core mass is $\approx 0.4M_\odot$, just below the SC limit. The core contracts until H shell burning is initiated at the core-envelope boundary, resulting in an increase in the surface temperature (this initial contraction produces the Henyey hook in the HR diagram). The He ash contributes to a growing core until the SC limit is reached. At this point the core contracts rapidly while the envelope expands (we will explore why this occurs below). This last contraction occurs on a thermal timescale, resulting in a Hertzsprung Gap.

We have made reference on several occasions to the concept of **shell burning**. This type of burning, in which a (usually quite thin) shell of the star is undergoing fusion in between the core and envelope, occurs at several evolutionary phases in the lifetime of stars. It is a new source of energy to power the star, and the lifetime of shell burning phases is $\approx 10 - 20\%$ of the main sequence lifetime for masses $< 2M_\odot$. Shell burning generally produces more luminosity than core burning due to the higher temperatures involved. The physics of shell burning gives rise to interesting behavior in the evolution of a star.

Stars that are undergoing shell burning usually have very extended envelopes. Furthermore the shell is usually very thin. Let's make the assumption that the overlaying envelope has negligible temperature and pressure, and that the shell has a thickness Δr_{sh} . Hydrostatic equilibrium demands that $dP/dr = \rho g$, which in the shell implies: $P_{\text{sh}}/\Delta r_{\text{sh}} \sim \rho_{\text{sh}} M_c/R_c^2$. This follows from the assumption that the core is providing the gravity. Invoking the ideal gas law in the shell leads to: $T_{\text{sh}}/\Delta r_{\text{sh}} \propto M_c/R_c^2$. Define $q \equiv \Delta r_{\text{sh}}/R_c$, so that $T_{\text{sh}} \propto q M_c/R_c$. For low-mass stars the He core is degenerate, which implies $P \propto \rho^{5/3}$ and $R_c \propto M_c^{-1/3}$, hence: $T_{\text{sh}} \propto q M_c^{4/3}$. These results have the following physical interpretation for low-mass stars: as

shell burning proceeds, the core grows in mass. Because the core is degenerate, it shrinks. The shell temperature increases, resulting in an increase in luminosity. The star climbs up the red giant branch, along the Hayashi line (because the envelope is convective).

Question: It was stated that shell burning is confined to a narrow region in the star (the shell is “thin”). This is a consequence of the very high temperature sensitivity of the burning rate (shell burning of H is via CNO). Why does a high temperature sensitivity give rise to a thin shell?

By considering simple models for the envelope (or running modern stellar evolution programs) one can connect the total luminosity to the *core* mass: $L/L_\odot \approx 2 \times 10^5 (M_c/M_\odot)^6$. Notice that the luminosity does *not* depend on the *total* mass. For this reason it is very difficult to determine the mass of a star on the RGB from its surface properties (we will see later that asteroseismology, by providing a window into the stellar interior, does offer a constraint on the mass of a star on the RGB).

One of the more remarkable facts about stars is that some are observed to be enormously larger than the Sun. In particular, giants and supergiants have radii of $100 - 1000 R_\odot$. Giants can either be cool ($T_{\text{eff}} \sim 4000 \text{ K}$ – close to the Hayashi line) or warm ($T_{\text{eff}} \sim 10,000 \text{ K}$). It is also somewhat remarkable that the physical origin of precisely why some stars become giants is still debated today (see e.g., Renzini et al. 1992, Iben 1993, Sugimoto et al. 2000, Stancliffe et al. 2009). In spite of the somewhat elusive origin, all stellar evolution codes have, since their earliest days, reliably produced giant stars.

One particularly simple explanation for why stars become red giants can be obtained by considering the virial theorem in the form $2K = -U$ where K and U are the kinetic and potential energy, respectively. If we assume that a change in structure occurs on a timescale shorter than the KH timescale, then by definition that means $K = \text{constant}$. By virial equilibrium we then also must have $U = \text{constant}$. For stars that display a clear core-envelope structure, we can write:

$$U \approx \frac{GM_c^2}{R_c} + \frac{GM_c M_{\text{env}}}{R_*} \quad (10.5)$$

where M_c and M_{env} are the core mass and envelope mass (both constant), and R_* is the total radius of the star (we have again assumed that the core provides the gravity). We can then compute:

$$\frac{dR_*}{dR_c} \approx - \left(\frac{M_c}{M_{\text{env}}} \right) \left(\frac{R_*}{R_c} \right)^2 \ll -1 \quad (10.6)$$

So for constant core and envelope masses, when the core contracts the stellar radius increases (significantly).

An even simpler way to see this is to write the potential energy as $U = -\alpha \frac{GM}{R_*}$. Again assuming that $K = \text{constant}$, as the core contracts α increases (where α is a constant that depends on the concentration of the mass), which means that R_* must increase to keep $U = \text{constant}$.

The assumption of $K = \text{constant}$ is not great, as the timescale of evolution is in fact comparable to or somewhat larger than t_{KH} . But the addition of some change in K will not in general change the *sign* of the effect above.

As with all arguments that rely on the virial theorem, it can only take us so far. The VT tells us that core contraction will produce envelope expansion (and visa versa, what Lamers & Levesque refer to as the ‘mirror principle’), but not *why*. In particular, it does not describe the microphysics that couples the core contraction to the envelope expansion. e.g., how does the envelope “know” to expand when the core is contracting? This is, I think, the origin of much of the head scratching on the question of why stars become giants.

We will conclude this lecture with a detailed discussion of the post-main sequence evolution of a $1M_{\odot}$ star with solar composition. This example is important not only because we happen to live near a $1M_{\odot}$ star, but also because such stars are long-lived so we see many examples of such stars in various evolutionary phases. For example, observations of globular clusters (GCs) have loomed large in our understanding of stellar structure. GCs are old, co-eval collections of stars, so the stars at the turnoff and all post-main sequence phases have masses of $\approx 0.8 - 1M_{\odot}$.

Figure 10.2 shows, in the right panel, the evolution in the HR diagram of a $1M_{\odot}$ star, from ZAMS through the end of nuclear burning. The left panel is known as a **Kippenhahn diagram** and it is a powerful, compact representation of the evolution of the interior structure of a star. The x-axis is time (often displayed in non-linear units) and the y-axis is the mass-coordinate. Gray regions mark where convection occurs, while hatched regions mark where nuclear burning occurs.

On the main sequence (points A-B) fusion is confined to the inner core, and only the outer $\lesssim 5\%$ of the star by mass is convective. As the star leaves the main sequence and transitions to the RGB (points B-C), a thin H burning shell develops, and the convective envelope penetrates much deeper into the interior. Notice that the location of the shell burning in mass

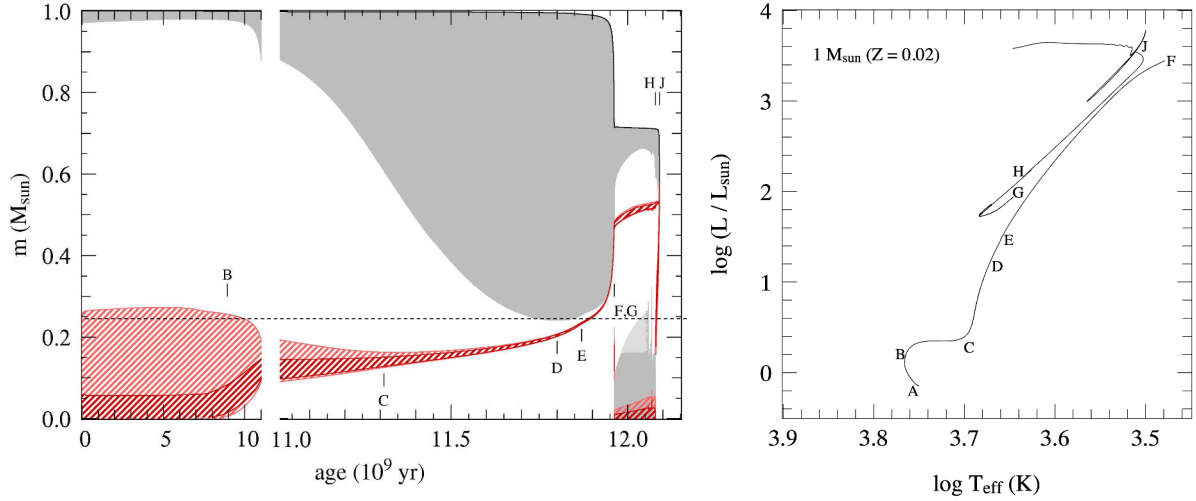


Figure 10.2: Evolution of a $1M_{\odot}$, solar metallicity star. Left panel is a Kippenhahn diagram, right panel shows an HR diagram. In the left panel, the internal structure is shown as a function of mass coordinate on the y-axis, and as a function of time on the x-axis. Dark gray regions are convective, while light gray areas are semi-convective. Red hatched regions show areas of nuclear energy generation for $\epsilon_{\text{nuc}} > 5L/M$ (dark red) and $\epsilon_{\text{nuc}} > L/M$ (light red). The horizontal black dashed line is to guide the eye in relation to the behavior at points D-E. From Pols (2011).

coordinates move outward over time (as it must, because the inner regions are increasingly He-rich). While not clearly depicted in the figure, the radius is expanding significantly during this time, for the reason we described earlier.

Something special happens at point D. Here the convective envelope reaches a point that, at earlier times, had experienced nuclear burning. This means that the chemical composition of the envelope will abruptly change, and since convective mixing is efficient, the surface layers of the star will reflect this change in composition. This process is known as the **first dredge-up** (there are two additional instances of dredge-up in later phases of evolution, and/or for higher-mass stars). The main change in the surface abundances is in the ratio of nitrogen to carbon – a reflection of CNO equilibrium driven by H burning. This is important because a change in the surface abundance is *observable*. Dredge-up provides a window into the state of the stellar interior.

Shortly after first dredge-up something else occurs. At point E the shell burning region reaches the point in the star where convection had earlier reached its deepest point. This means that the shell will experience an abrupt change in the mean molecular weight, which results in an abrupt change (decrease) in the emergent luminosity. As a consequence, the

star transverses the same luminosity three times (up, down, and back up again). This is known as the **luminosity bump** because, in a population of stars (e.g., a star cluster), there will be a bump in the luminosity function along the RGB owing to the excess number of stars at the luminosity corresponding to point E.

While not explicitly depicted in the figure, the RGB phase comes to an end when the core temperature and density rise enough to ignite He in the core (point F). For low-mass stars ($\lesssim 2M_{\odot}$) this is an explosive process because the core is degenerate (which means the temperature is uncoupled from the density and pressure). This is known as the **helium flash**. We will return to this and later phases of evolution in subsequent lectures.

Summary: In this lecture we introduced the concept of the core-envelope structure within stars. This is an important conceptual tool for understanding stellar interiors, and is widely applicable. At the core-envelope interface there is sometimes a thin shell of material undergoing nuclear burning. In such cases we have a core-shell-envelope structure. In this lecture we discussed the core mass required to support the mass of the overlaying envelope (the Schönberg–Chandrasekhar Limit), the nature of shell burning, and why stars inflate to red giant dimensions. These principles are general and will be applied throughout the next several lectures to understand the basic principles shaping the evolution of low and high mass stars.

A common theme will be the striking complexity of post-main sequence evolution. The complexity is driven by the many options available to the state of the stellar interior: the core may be radiative or convective, stabilized by ideal gas pressure or degeneracy pressure, may be isothermal or have a temperature gradient set by nuclear energy production. There can be zero, one, two or more distinct shell burning regions. The inter-shell and envelope regions may be radiative or convective. Of course, the computer can solve the equations and give us the “answer”, but we are interested in extracting and understanding the basic principles governing the interior structure and observed properties of stars.

Chapter 11

Stellar Winds & Mass Loss

Stellar mass loss plays a very important role in the evolution of stars. We observe, for example, that stars like the Sun end up as $\sim 0.5M_{\odot}$ white dwarfs. How is all that mass lost, and in which evolutionary phases? Mass-loss also plays an important role in the evolution of massive stars, producing the distinctive Wolf-Rayet stellar type, and will affect what kind of supernova light curve is observed. Mass-loss rates vary widely, from $\approx 10^{-14}M_{\odot} \text{ yr}^{-1}$ for the Sun, to values as high as $\approx 10^{-4}M_{\odot} \text{ yr}^{-1}$ for dust-driven winds from AGB stars. In this lecture we will discuss some of the theory and phenomenology of stellar winds.

At a basic level, a stellar wind requires a source of momentum in order to accelerate the gas above the escape speed, where $v_{\text{esc}} = \sqrt{2GM/r}$. There are three broad sources of momentum: the radiation field, thermal winds (driven by gradients in the gas pressure), and waves (sound waves, shocks, Alfvén waves, etc). Following Lamers & Levesque (Ch 15.1) there are four broad categories of stellar winds.

- **Line-driven winds** are driven by radiation pressure on ionized atoms (including C, N, O, and Fe-group elements). This mechanism is important for hot stars (where the ionization of these elements is significant), and can generate fast winds with terminal velocities of $\sim 10^3 \text{ km s}^{-1}$.
- **Dust-driven winds** are driven by radiation pressure on dust grains. This mechanism is important for cool extended stars such as red supergiants and AGB stars, and produces slow winds with velocities of $\sim 10 \text{ km s}^{-1}$. Pulsation is believed to be an important component to this model, as otherwise the dust-forming region is at a density that is too low to be effective at driving a wind.

- **Coronal winds** are driven by gas pressure sourced by the high temperature of the stellar corona (which can be as hot as 10^6 K). This mechanism is responsible for the solar wind, and perhaps the winds of cool main sequence stars and giants.
- **Alfven wave-driven winds** are driven by magnetic waves. The waves can generate an outward pressure gradient that accelerates ionized gas.

Independent of the driving mechanism, we often find ourselves considering a stationary wind, in which there is no time-dependent phenomena. In this case we have $\dot{M} = \text{constant}$ and $v = \text{constant}$. The hydrodynamic equation of continuity also implies $\nabla \cdot (\rho v) = 0$. For spherical symmetry this equation implies $\dot{M} = 4\pi r^2 \rho v$, and so a stationary spherical wind implies a density profile in the wind of the form: $\rho \propto r^{-2}$.

Let's consider the case of line-driven winds in some detail. In hot stars the most abundant elements will be multiply ionized. Moreover, many of the strongest transitions will be in the UV, where hot stars also emit most of their light. These ions will be accelerated by momentum transfer from the radiation field during the process of absorption or scattering. While these ions represent a small fraction of all ions, their frequent interactions via collisions with other ions (and free electrons) will drag the rest of the gas along.

We can make a crude estimate of the mass-loss rate due to line-driven winds as follows. Assume a particular ion has a very strong absorption line at a wavelength λ_0 that corresponds to the peak of the Planck function for the star. Assume that the line is optically-thick and absorbs or scatters all photons at λ_0 . Furthermore, assume that the gas velocity increases from $v = 0$ at the photosphere to v_∞ at large distances. Due to the Doppler shift, all photons in the frequency range ν_0 to $\nu_0(1 + v_\infty/c)$ are absorbed or scattered. The momentum of a photon is $E/c = h\nu/c$ and the total momentum of all photons leaving the star per second is L/c . If all photons were absorbed, then the total momentum flux would be L/c – instead we need to compute the number of photons that will experience absorption due to this single line:

$$L_{\text{abs}}/c = \int_{\nu_0}^{\nu_0(1+v_\infty/c)} 4\pi R_*^2 F_\nu/c d\nu = L_{\nu,0} \Delta\nu/c \quad (11.1)$$

where $L_{\nu,0}$ is the luminosity at ν_0 and $\Delta\nu = \nu_0 v_\infty/c$.

The total momentum loss rate due to the wind is $\dot{M}v_\infty$, while the momentum transfer rate due to absorption/scattering is L_{abs}/c . In equilibrium these two will be equal:

$$\dot{M}v_\infty = L_{\text{abs}}/c = L_{\nu,0}\nu_0 v_\infty/c^2 \quad (11.2)$$

At the peak of the Planck function $L_{\nu,0}\nu_0 \approx 0.6L$. If we round $0.6 \sim 1$ then we find that one strong line can drive a mass-loss rate of $\dot{M} \approx \frac{L}{c^2}$. For a hot massive star, $L \approx 10^6 L_\odot$, so $\dot{M} \approx 4 \times 10^{-8} M_\odot \text{ yr}^{-1}$ per spectral line. If the spectrum contains ~ 100 strong spectral lines, then the total mass-loss rate would be $4 \times 10^{-6} M_\odot \text{ yr}^{-1}$ which is in the ballpark of observed estimates of mass-loss rates of hot stars.

Note that the maximum possible mass-loss rate of radiation-driven winds can be estimated by setting the momentum transfer rate to L/c , which yields: $\dot{M}_{\text{max}} = \frac{L}{v_\infty c}$. For a typical hot main sequence star, $v_\infty \approx 2v_{\text{esc}} \approx 2000 \text{ km s}^{-1}$, so the maximum mass-loss rate of a $10^6 L_\odot$ star is $\approx 10^{-5} M_\odot \text{ yr}^{-1}$.

Finally, note that in the above calculations we have assumed that each photon can be “used” (absorbed/scattered) only once. In reality, there are many spectral lines, and so photons can be absorbed/scattered multiple times, resulting in higher mass-loss rates than we have estimated.

Strong line-driven stellar winds are ubiquitous among luminous hot stars, including O, B, and WR stars. Wind velocities can reach 10^3 km s^{-1} and mass-loss rates range from $10^{-10} - 10^{-5} M_\odot \text{ yr}^{-1}$. These estimates imply that stars on the main sequence with $M \gtrsim 30 M_\odot$ will lose a significant fraction of their mass ($\approx 5 - 30\%$) while on the main sequence.

Measuring mass-loss rates is challenging and ultimately relies on radiative transfer models. A striking observational feature indicating an outflow of material is P Cygni line profiles (see the accompanying slides for details).

Dust often plays an outsized role in astronomy. By mass it is almost always insignificant, because dust is made of metals, and metals are rare (dust-to-gas ratios are usually in the range 10^{-2}). However, dust has a very high opacity, and so plays an important role in radiation transfer when it is present.

Typical dust grains have a mean (internal) density of $\rho_d \approx 1 \text{ g cm}^{-3}$, and they have an average size of $a \approx 0.05 \mu\text{m}$ (with a wide, usually power-law distribution of sizes). This implies a total absorption coefficient of $\kappa_d = \rho_d \pi a^2 Q \approx 10 \text{ cm}^2 \text{ g}^{-1}$, for an extinction efficiency factor of $Q = 10^{-2}$. The quantity Q encapsulates the optical properties of the grain and is in general a function of grain size, grain composition, and wavelength. If we require that the radiation pressure force exceed the force of gravity, then we arrive at the following condition:

$$\frac{L\kappa_d}{4\pi r^2 c} > \frac{GM}{r^2} \quad (11.3)$$

which implies a minimum luminosity for dust-driven winds of $L/L_\odot > 10^3 M/M_\odot$. Luminous RGB stars and red supergiants have luminosities well in excess of this requirement.

Circumstellar dust will only form at relatively low temperatures – typical condensation temperatures for grains are in the range of 1100 – 1500 K. The temperature of the dust depends on the equilibrium between the energy of the absorbed radiation and their subsequent cooling radiation. The energy absorbed per second by a spherical grain of radius a at a distance r from a star of luminosity L and temperature T_{eff} is:

$$e_{\text{in}} \approx \pi a^2 Q_{\text{abs}} \frac{L}{4\pi r^2} = \pi a^2 Q_{\text{abs}} \sigma T_{\text{eff}}^4 \frac{R^2}{r^2} \quad (11.4)$$

where Q_{abs} is again the efficiency factor for absorption by the grain particle. We have ignored the geometrical factor associated with the fact that the flux does not fall off quite as rapidly as $(R/r)^2$ near the surface. The energy radiated by a grain of temperature T_d per second is:

$$e_{\text{out}} = 4\pi a^2 Q_{\text{em}} \sigma T_d^4 \quad (11.5)$$

setting $e_{\text{in}} = e_{\text{out}}$ and $Q_{\text{em}} \approx Q_{\text{abs}}$ yields: $T_d^4 \approx T_{\text{eff}}^4 \frac{1}{4} (R/r)^2$. For a typical cool (super)giant with $T_{\text{eff}} \approx 2500$ K, dust will form at distances of about $2R$.

But there is a problem. Absent a wind, the density at $2R$ is very low, so that dust will not form. But if there is no dust, then there will be no wind! Note that for the very coolest stars with $T_{\text{eff}} < 2000$ K this is less of a concern, as dust can form closer to the photosphere.

At this point one appeals to pulsating stars. Observations indicate that many extended giants have large pulsation amplitudes (> 1 mag) with periods of > 100 d. One of the effects of pulsations is to elevate material to larger radii, thereby giving a higher density of material than would be expected from a hydrostatic atmosphere. In short, the density at $2R$ can be much higher, allowing dust to form and a wind to develop. The details are very complicated, and 3D hydrodynamic simulations with radiative transfer and dust grain formation are required. Such simulations are demanding.

OH/IR stars are very cool, evolved AGB and red supergiant stars that have pulsation periods > 1000 d, wind velocities of ~ 10 km s $^{-1}$, and mass-loss rates of $\approx 10^{-4} M_\odot$ yr $^{-1}$. These stars are IR-bright and enshrouded in dust. They are the most extreme versions of mass-losing cool stars. The mass loss rates are determined by comparison to detailed radiative transfer models that include grains and the density structure in the wind.

A complete theory of stellar winds and mass loss does not exist. And yet, various aspects of stellar evolution are quite sensitive to mass loss, as we will see in later lectures. For this

reason phenomenological models of mass-loss are employed in stellar evolution calculations. These models are either based on observations or more detailed wind calculations. Specific models are used for hot main sequence stars, Wolf-Rayet stars, cool stars, etc.

One of the most popular mass-loss prescriptions is the **Reimers relation**:

$$\dot{M} = 4 \times 10^{-13} \eta_R \frac{L}{L_\odot} \frac{R}{R_\odot} \frac{M_\odot}{M} M_\odot \text{ yr}^{-1} \quad (11.6)$$

which was presented in Reimers (1975) based on observations of *six* cool giants. The fudge-factor η_R is often varied to produce better agreement with observations. This scheme is almost universally used to model mass-loss along the RGB. A cynic might observe that this is simply the combination of fundamental variables that produces a quantity with units of mass per time. Either way, it seems to work reasonably well at describing observations of mass-loss for cool stars.

Finally, it's important to recognize that non-steady state mass-loss likely occurs in certain special phases of stellar evolution and may be catastrophic to the star. For example, eruptive mass-loss produces some of the most spectacular observed phenomena, including η Carinae, which underwent a great eruption occurring from 1837–1843 that resulted in a mass loss of $10 - 20M_\odot$! The origin of these eruptions is completely unknown; ideas include physics related to stellar binarity, or the star exceeding the Eddington limit. The impact of these eruptive phases on stellar evolution is not well-understood nor even well-explored from a theoretical point of view.

Question: Based on this lecture, what can you surmise about the metallicity-dependence of mass-loss for different stellar types? How might this dependence affect the evolution of low and high-mass stars?

Chapter 12

Evolution of Low-Mass Stars

In this lecture we will continue to follow the evolution of low-mass stars ($M \lesssim 8M_{\odot}$) until they exhaust all sources of nuclear burning. You will notice that this lecture is light on equations. While some of the phenomena discussed here can be estimated via back-of-the-envelope calculations, most of the interesting behavior requires realistic stellar evolution calculations. The discussion here is based on such calculations.

In earlier lectures we learned that stars burn hydrogen on the main sequence until the source of fuel is exhausted. At this point the core contracts until a shell of hydrogen begins burning. As the core contracts the envelope expands, resulting in a giant star. The star begins its ascent up the red giant branch (RGB). See Figure 12.1 for the key phases of evolution for $1M_{\odot}$ and $5M_{\odot}$ stars.

During the RGB the core is inert but is setting the gravitational potential and, via the ideal gas law, the temperature of the shell. The temperature of the shell sets the luminosity of the shell, and, since shell burning is the only source of energy during the RGB, the total luminosity of the star. The core therefore controls much of the evolutionary behavior along the RGB. As the shell consumes hydrogen it grows the helium core, resulting in higher shell temperatures and hence higher luminosities. See Lecture 10 for details on the behavior of stars along the RGB.

Models tell us that when the core reaches a mass of $M_c \approx 0.5M_{\odot}$ (which occurs when $L \sim 10^3 L_{\odot}$), the core temperature is $T_c \sim 10^8$ K. At this temperature helium fusion begins.

The onset of core helium burning is very different for stars below and above $\approx 2M_{\odot}$. When $M > 2M_{\odot}$, the core is non-degenerate, and helium fusion begins “gently”, in a manner

analogous to the onset of hydrogen burning on the main sequence. By “gentle” we mean that the pressure increases as the temperature increases (via the ideal gas equation of state), which leads to an expansion of the core and a decrease in the temperature. The process is self-regulating and leads to an equilibrium state. Compare the RGBTip-ZACHeB transition in the right panels of Figure 12.1.

Below $\approx 2M_{\odot}$ the helium core is degenerate, which entails dramatically different behavior (the exact threshold depends on various factors and is closer to $1.75M_{\odot}$ in modern models. Be careful: this threshold is different than the radiative-convective transition on the main sequence, which occurs at $\approx 1.5M_{\odot}$). In this case, temperature and pressure are uncoupled, so that an increase in temperature does not lead to a change in the hydrostatic structure of the star: the star is not self-regulating. This leads to a runaway process until the central temperature increases so much that degeneracy pressure is lifted (i.e., the star moves onto the ideal gas relation of $T_c - \rho_c$). This runaway process is called the **helium flash**. There are several fascinating aspects of the flash. Neutrino cooling is sufficiently vigorous so that the central temperature is slightly cooler than the off-center core. This means that the flash occurs off-center. The luminosity produced by the helium flash is enormous: $L_{\text{He}} \sim 10^{10}L_{\odot}$! This luminosity is sustained for a very brief period of time, and is never observed, because most of the energy goes into lifting degeneracy in the core. Finally, the flash is actually multiple flashes as the degeneracy in the core is lifted and steady-state burning is established. You will explore the helium flash in detail in HW3.

Question: Estimate the maximum time that a star can sustain the luminosity produced in the He-flash assuming the core is (and remains) in virial equilibrium.

One way or another all stars eventually settle down into a state of steady core helium burning in what is sometimes referred to as the **helium main sequence**. We can rely on several of the arguments from previous lectures regarding the hydrogen main sequence to understand the behavior of stars on the helium main sequence. The two main differences are the different mean molecular weight and the distinction between the stellar core (relevant for helium burning) and the overall initial mass (relevant for hydrogen burning). Recall from Lecture 8 that luminosity is proportional to μ^4 or $\mu^{7.5}$ at fixed mass (depending on the form of energy transfer). The mean molecular weight for a solar metallicity mixture ($Z = 0.02$, $Y = 0.27$) is $\mu = 0.61$, while for a helium-rich core it is $\mu \approx 1.33$. So the helium main sequence will be $(1.33/0.61)^4 \approx 22$ times brighter than the hydrogen main sequence. This higher luminosity combined with the smaller amount of energy liberated per nuclear reaction (7.3 vs. 22 MeV), results in a substantially shorter core helium burning lifetime compared

to the hydrogen main sequence.

Consideration of the core mass results in some additional complications. Again, for stars $M > 2M_{\odot}$, where helium fusion occurs “gently”, the behavior will largely mimic that of the hydrogen main sequence, in which one can derive relations of the form $L \propto M^{\beta}$ depending on the form of energy transfer and the sources of opacity. For $M < 2M_{\odot}$ the degenerate core implies that the central temperature depends only on the core mass. This means that the helium flash occurs at approximately the same core mass, $M_c \approx 0.5M_{\odot}$, *independent of initial mass*. This has two important implications. First, the maximum luminosity achieved on the RGB will be approximately constant for initial mass below $2M_{\odot}$ and hence with ages greater than the turnoff age of a $2M_{\odot}$ star, i.e., $\approx 10^9$ yr. In other words, the so-called **tip of the red giant branch** (TRGB) is quite stable with age, and furthermore is understood on basic physical grounds. It is for these reasons that the TRGB is often used as a standard candle. The primary variable affecting the luminosity of the TRGB is the metallicity (composition), which unfortunately is often not well-known in many observational contexts.

The second important implication of the constant core mass at which the helium flash occurs is the following: the luminosity during core helium burning is approximately constant for all such stars that had degenerate cores when ascending the RGB. The luminosity is constant at $\sim 10^2 L_{\odot}$. Stars in this phase of evolution (core helium burning with initial mass $< 2M_{\odot}$) are known as **red clump** (RC) or **horizontal branch** (HB) stars. These terms are derived from their appearance in the HR diagram. Specifically, metal-rich populations display a clear clump of stars at a constant luminosity that is near, but offset blueward (toward hotter temperatures) from the RGB. At lower metallicities one more frequently sees a horizontal locus of stars at constant luminosity and varying T_{eff} – this is known as the HB.

Stars on the HB/RC have the following interior structure (moving from the core to the surface): an inner core in which helium fusion occurs. Energy transfer in this region is convective because of the high flux. Beyond the burning region is an inert helium core. Beyond the helium core is a region of hydrogen shell burning. The shell has luminosity comparable to the helium fusion in the core. Beyond the shell burning is an inert hydrogen envelope. Since the initial core mass is fixed, T_{eff} depends on the envelope mass. Hence T_{eff} on the HB depends on both the initial mass and stellar mass-loss. At fixed age the turnoff mass depends on metallicity such that lower metallicity populations have lower turnoff masses. In the most extreme cases the temperature along the HB can reach 20,000 K.

As helium burning proceeds, a carbon-oxygen (CO) core builds up. The basic reactions are triple- α , i.e., $3 \text{ He} \rightarrow \text{C}$, and helium+carbon burning, i.e., $\text{He} + \text{C} \rightarrow \text{O}$. Stars spend approx-

imately ≈ 100 Myr in the core helium burning phase for $M < 2M_{\odot}$. At higher masses the lifetime is approximately 10% of the main sequence lifetime.

After core helium burning, the star ascends the **asymptotic giant branch** (AGB), so-named because the track asymptotically approaches the RGB at high luminosities. During the AGB phase the interior structure is: a) inert, degenerate CO core; b) He fusion shell; c) He-rich inter-shell zone; d) H fusion shell; e) H-rich convective envelope. As in the RGB phase, the AGB luminosity is mostly dictated by the core luminosity: $L/L_{\odot} \approx 6 \times 10^4 (M_c/M_{\odot} - 0.522)$. The AGB phase lasts for $\approx 10^7$ yr.

Lots of interesting things happen during the AGB phase (see Figures 12.1 and 12.2 below):

- **2nd dredge-up (2DU)**: This is similar to the 1DU discussed in Lecture 10. The 2DU requires extinction of the H shell burning, which occurs only for $> 4M_{\odot}$. Lower mass stars do not undergo 2DU.
- **Thermal pulses (TPs)**: During most of the evolution along the AGB the bulk of the luminosity is provided by the H shell (this is the early-AGB, or EAGB phase). As He ash accumulates, He shell burning eventually occurs in an unstable manner (similar in some respects to the He flash; see the end of these lecture notes for a discussion of the shell burning instability). Stars are expected to undergo many TPs in this phase (known as the TP-AGB phase), ranging from $\sim 5 - 30$. The overall lifetime of the TP-AGB phase is expected to be $\lesssim 10^6$ yr. The intervals between pulses can be $10^3 - 10^5$ yr depending on mass. The duration of a pulse is $\sim 10^2$ yr – short by stellar evolution standards, but long enough such that we have not directly observed a star undergoing a thermal pulse. See Fig. 18.3 in Lamers & Levesque for a schematic diagram of the TP-AGB phase and Figure 12.2 below for a more detailed view.
- **3rd dredge-up (3DU)**: An additional dredge-up process can occur during the TPs that can substantially change the surface composition of a star. It is during this phase that **carbon stars** can be produced, in which the surface composition has $C/O > 1$. When $C/O > 1$ the molecular equilibrium on the surface undergoes a phase transition, in which carbon-bearing molecules (including C_2 , CN, CO, HCN, C_2H_2 , and C_3) dominate the appearance of the star.
- During the AGB phase, and especially the TP-AGB phase, mass-loss rates can reach $10^{-4} M_{\odot} \text{ yr}^{-1}$ (sometimes referred to as the superwind phase). This is the phase during

which stars eventually shed the remainder of their envelopes, and end the luminous phase of their lives.

The AGB phase concludes with the removal of the envelope via mass-loss. When this occurs the star makes a rapid excursion to hotter T_{eff} at approximately constant L_{bol} . T_{eff} can reach values of 100,000 K as the remaining very thin envelope is consumed. This excursion is very rapid, occurring on a timescale of $10^2 - 10^3$ yr. The timescale is set by how rapidly the remaining envelope is consumed, whether by burning or mass-loss (or some combination thereof – the details are uncertain). This is known as the post-AGB phase. When the star reaches $T_{\text{eff}} \approx 20,000$ K the central star is hot enough to ionize the stellar ejecta surrounding it. This results in the spectacular **planetary nebula** phenomenon. After the star reaches its maximum temperature it cools to the white dwarf stellar graveyard (we will return to this in later lectures).

We will now briefly discuss the **shell burning instability** relevant for the production of thermal pulses. Consider a thin shell of width D . If the shell expands at constant mass ($m \sim \rho r_0^2 D = \text{constant}$) then

$$\frac{d\rho}{\rho} = -\frac{dD}{D} \quad (12.1)$$

Now imagine excess heat in the shell (e.g., due to an increase in ϵ_n). The shell will want to expand (e.g., as an ideal gas), i.e., $dD/D > 0$ and hence $d\rho/\rho < 0$. However, if the shell is thin ($D/r \ll 1$) then $dr/r \ll 1$, which means that the layers above the shell are hardly lifted, so the weight remains constant and hence $dP/P \approx 0$. The ideal gas law demands:

$$d\rho/\rho = dP/P - dT/T \quad (12.2)$$

hence if $dP/P \approx 0$ then $d\rho/\rho = -dT/T$. But we argued above that $d\rho/\rho < 0$, so we must have $dT/T > 0$! This is clearly an unstable situation – adding heat results in a hotter shell. This is known as the **thin shell instability**, and is the physical mechanism behind the thermal pulses in TP-AGB stars, and also of classical novae. The instability is quenched either when the fuel source is consumed or when the increase in pressure is so great that $dP/P \approx 0$ is no longer a good assumption.

We have now charted out the evolution of low-mass ($M \lesssim 8M_{\odot}$) stars from pre-main sequence contraction through the formation of a hot white dwarf. The evolution of the white dwarf will be picked up in a later lecture. It is important to understand the major evolutionary phases

and their key characteristics, both in terms of surface properties (i.e., the HR diagram) and internal structure. It is also important to understand the key mass dependences and general metallicity-dependence of these various phases.

The boundary between “low” and “high” mass evolution is set by whether the star can continue burning beyond helium in its core. The mass boundary depends on metallicity and to some extent on uncertain stellar physics, but is generally expected to be in the range of $8 - 10M_{\odot}$.

At the lowest masses ($\lesssim 0.6M_{\odot}$, depending on metallicity) the main sequence lifetime is much longer than the age of the universe, so there is usually no need to consider post-main sequence evolution of such objects.

Question: What might our distant ancestors see if they could live long enough to observe the post-main sequence evolution of a $0.4M_{\odot}$ star? What would be the final fate of such an object?

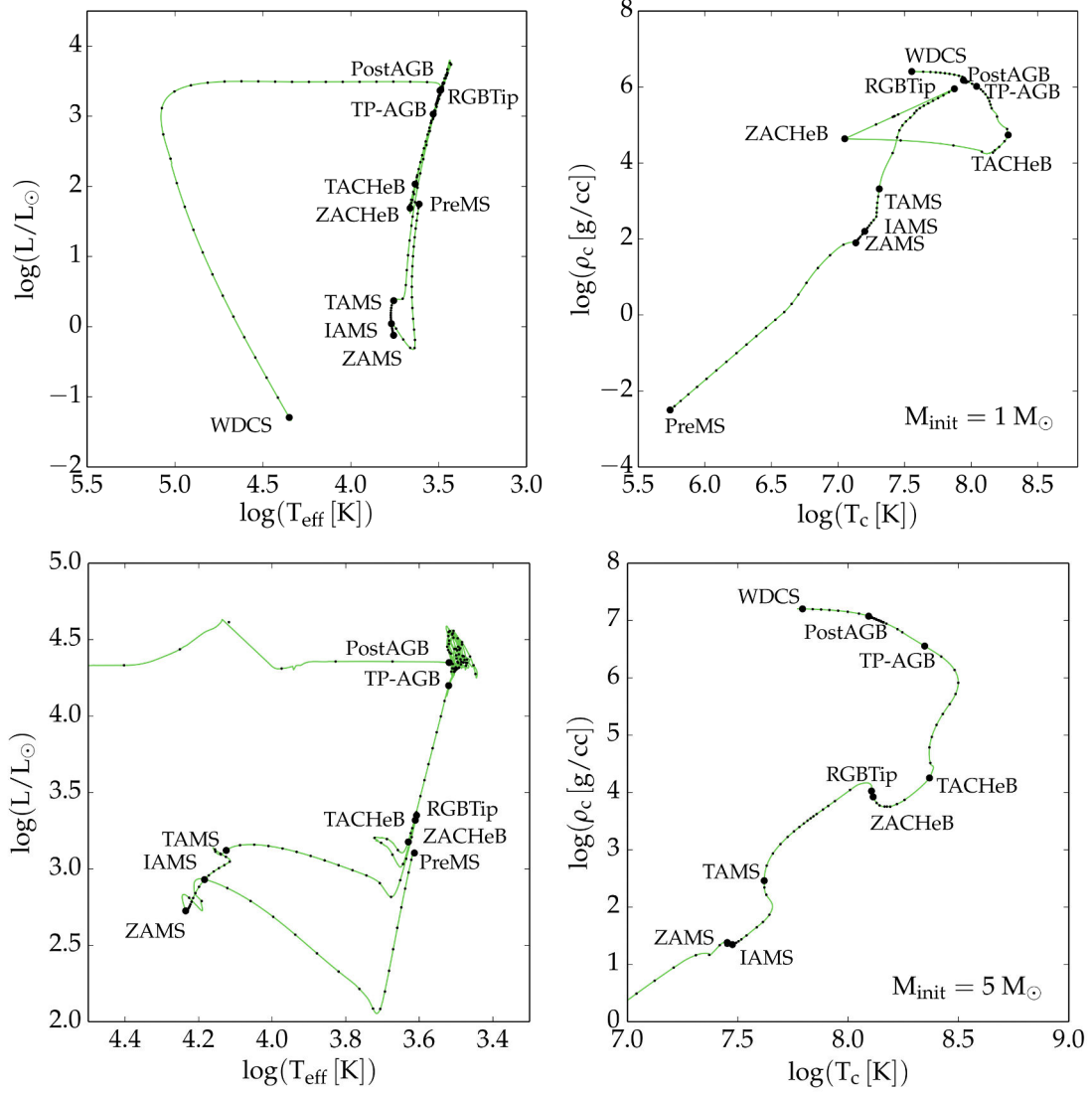


Figure 12.1: Evolution in the HR diagram (left panels) and the central density - central temperature plane (right panels) for $1 M_{\odot}$ (top panels) and $5 M_{\odot}$ (bottom panels) solar-metallicity stars. Important phases are marked: pre-main sequence (PMS), zero-age main sequence (ZAMS), intermediate-age main sequence (IAMS), terminal-age main sequence (TAMS, when central H is exhausted), zero-age core He burning (ZACHeB), terminal-age core He burning (TACHeB), the TP-AGB phase, post-AGB phase, and white dwarf cooling sequence (WDCS). Notice the very different behavior in the right panels around the RGB tip for the 1 and $5 M_{\odot}$ models. The small black points are *not* uniformly spaced in time. Reproduced from Dotter 2016.

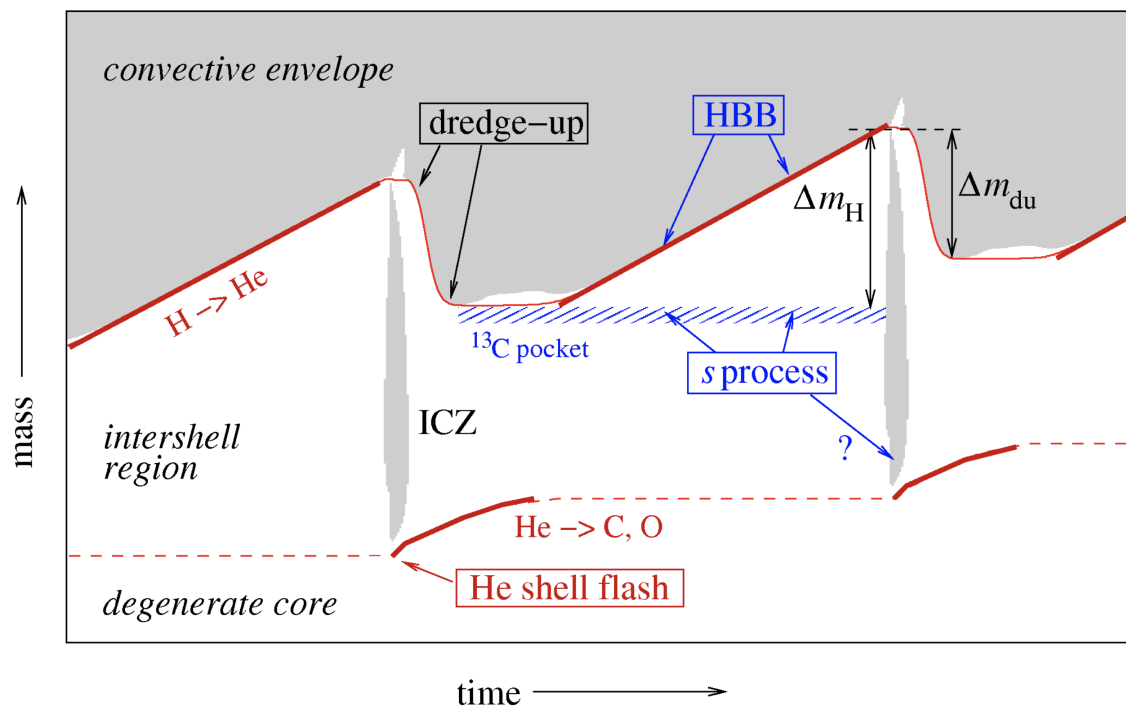


Figure 12.2: Schematic Kippenhahn diagram of a star experiencing two AGB pulses. Grey regions mark convective zones. “ICZ” is the inter-shell convective zone, driven by the He-shell flash. Dredge-up events are marked, as is the hot-bottom burning region, and the site of the s-process. The axes are highly non-linear. Reproduced from Pols 2011 (their Fig. 11.3).

Chapter 13

Evolution of Massive Stars

Late-stage evolution of massive stars ($> 8M_{\odot}$) is the most uncertain regime of stellar evolution theory. This uncertainty is primary due to the uncertain effects of stellar mass-loss, rotation, and binary stellar evolution. Stars in this phase can also become unstable, which may lead to dramatic eruptive behavior (e.g., the luminous blue variable [LBV] phase) that cannot be adequately captured in 1D models. The challenge imposed by binarity and rotation are twofold: 1) they reflect initial conditions of the system, 2) they are fundamentally 3D processes that can only be crudely captured in a 1D model. The fact that stars can have different initial rotation speeds and binary star parameters (e.g., mass-ratio and orbital separation) means that much larger grids of models must be computed, which is a computational burden. For lower mass stars mass-loss becomes very large only at the end of the AGB phase, and so has a minor effect on the evolution of stars with $M < 8M_{\odot}$. However, massive stars can lose substantial fraction of their initial mass while still on the main sequence, and so the entire evolutionary history of massive stars can be affected by stellar mass-loss. We will cover binary stellar evolution in later lectures.

Massive stars play a very important role throughout astronomy: they can explode as core-collapse SNe, leaving behind neutron stars or black holes (and are therefore important sources for gravitational wave events); they enrich the interstellar medium with metals and therefore play a large role in shaping the chemical evolution of galaxies; they inject large quantities of energy and momentum into the surrounding gas, which has a large impact on the thermodynamic state of the galaxy.

There is an interesting transition region between low and high mass stars, which is roughly the range $8 \lesssim M \lesssim 10M_{\odot}$. In this regime (known as the super-AGB phase) the CO core is able to undergo some burning to produce an O+Ne+Mg degenerate core. Such a core,

if sufficiently massive (and hence dense) will undergo electron capture reactions (mostly on ^{24}Mg), which removes pressure and can lead to the collapse of the core. An electron capture occurs when a nucleus absorbs an electron and the proton and electron combine to create a neutron (and an electron neutrino). Such events are called electron-capture supernovae, and they leave no remnant behind. The details of this phase are uncertain and depend on metallicity.

In spite of all the complexity associated with massive stars, we can identify a few basic features: 1) degeneracy pressure is never the dominant source of pressure support in the core. For this reason burning can proceed through iron. 2) The core will eventually collapse, leading in most cases to a supernovae. 3) Each burning phase occurs on shorter timescales, because the energy liberated per nuclear reaction ($\propto \Delta mc^2$) is smaller for heavier ion fusion, and because neutrino cooling results in substantial loss of energy from the core. Phases after carbon burning last less than a year. 4) Massive stars develop an “onion skin”-like structure in which the most massive elements are concentrated toward the center.

The behavior of massive stars is different above and below a threshold mass of $\approx 25M_\odot$. Below this mass, massive stars are believed to end their lives as red supergiants (along the Hayashi line). While in the red supergiant phase, some stars undergo **blue loops** during core He burning. The blue loop extends to temperatures of $T \gtrsim 10^4$ K. Whether or not a star will undergo a blue loop depends sensitively on the gravitational potential of the core as well as details of convective core overshooting and metallicity. This phase is another example of a “fact” from stellar evolution that is very difficult to explain in terms of basic physics (see Walmswell et al. 2015). The blue loops intersect the cepheid instability strip in the HR diagram. Cepheid variables are therefore massive stars in the blue loop phase.

For masses $\gtrsim 25M_\odot$ post-main sequence evolution is strongly affected by mass-loss. In the $25 \lesssim M \lesssim 50M_\odot$ range, stars evolve to the red supergiant phase, where mass-loss eventually strips the envelope and the star evolves toward hotter temperatures where they appear as **Wolf-Rayet (WR) stars**. For $M \gtrsim 50M_\odot$, mass-loss is so high on the main sequence that they do not even reach the red supergiant phase; they evolve off the main sequence and immediately evolve into WR stars.

WR stars are very luminous ($L > 10^5 L_\odot$) and very hot, with $T_{\text{eff}} > 30,000$ K. These stars have strong broad emission lines, which is indicative of strong stellar winds. There are several sub-categories of WR stars, including WN (nitrogen-rich) and WC (carbon-rich). These stars have had their hydrogen envelopes stripped off (presumably by mass-loss), thus exposing the nuclear processed material deeper in the star.

Advanced Burning. Each subsequent burning stage proceeds on a significantly shorter timescale. This is governed mostly by the increasing importance of energy loss due to neutrinos (due to the neutrino cooling process discussed in Lecture 5), and also by the decreasing amount of energy liberated per nuclear reaction. The general progression of core fusion is broadly: hydrogen, helium, carbon, neon, oxygen, and finally silicon. For a $15M_{\odot}$ star, core hydrogen burning lasts ≈ 11 Myr, while silicon burning lasts 18 d! For a $25M_{\odot}$ star these timescales are ≈ 7 Myr and 0.7d.

Carbon burning proceeds via the $^{12}\text{C}+^{12}\text{C}$ reaction, which produces a mixture of products, mainly ^{20}Ne and ^{24}Mg . Most of the energy produced escapes in the form of neutrinos. When core carbon burning finishes and before the next core burning phase begins, shell carbon burning occurs. As in earlier burning phases, this structure develops because the carbon-exhausted core must contract. When it does so, conditions in the shell region are conducive to carbon burning before the core becomes conducive to the next burning stage.

When the core temperature reaches 1.5×10^9 K, neon is burned into oxygen and magnesium by a combination of photodisintegration and α -capture. At this temperature a sufficient number of photons have energies in the MeV range that they can break up the relatively fragile ^{20}Ne atom via: $^{20}\text{Ne} + \gamma \rightarrow ^{16}\text{O} + ^4\text{He}$. Oxygen then immediately undergoes an α -capture reaction to produce ^{24}Mg . The first reaction is endothermic, but the second is exothermic, and the total effect is a net production of energy. The duration of neon burning is only ~ 1 yr. At slightly higher core temperatures ($\approx 2 \times 10^9$ K) oxygen burning begins via the $^{16}\text{O}+^{16}\text{O}$ reaction, which produces mainly ^{28}Si and ^{32}S . The core at this point has a typical mass of $\approx 1.0M_{\odot}$. Oxygen burning also takes about a year, in spite of the even higher neutrino losses compared to neon burning. This is because of the higher overall oxygen mass fraction, and the larger energy gain from oxygen burning. As with carbon burning, after core oxygen burning one or more oxygen burning shells develop before the next core fusion phase begins. At this point, the evolution is so rapid ($\lesssim 1$ yr) that the overlying layers (e.g., the hydrogen/helium envelope, and any shells of helium or carbon burning) are frozen into the stellar structure.

The final phase before core-collapse is known as “silicon burning”. This is not a single reaction, but is instead a complex set of photodisintegration and α -capture reactions that take place simultaneously. The end state can be described by the nuclear equivalent of the Saha equation and for $T > 4 \times 10^9$ K the final mixture is close to **nuclear statistical equilibrium (NSE)**, in which the most abundant element is the one with the highest binding energy, i.e., the iron group elements, most prominently ^{56}Fe .

Stellar Rotation. Massive stars are observed to have high equatorial rotation velocities, ranging from $200 - 400 \text{ km s}^{-1}$. These high rotation rates are likely inherited from the collapse of the protostar. Unlike the case of low-mass stars, in which magnetic braking efficiently removes angular momentum, in massive stars without convective envelopes, there is no efficient braking mechanism. As we will see later, binary mass transfer can also spin up a star. Rotation is therefore expected to play an important role in the evolution of massive stars.

A key concept is the **critical velocity**, which we can define as the velocity at which the centrifugal acceleration is equal to the gravitational acceleration. If we define an effective gravity as $g_{\text{eff}} = g_{\text{grav}} - g_{\text{centrifugal}}$, then the critical velocity occurs when $g_{\text{eff}} = 0$. If $g_{\text{grav}} = GM/R^2$ and $g_{\text{centrifugal}} = V^2/R$ then $V_{\text{crit}} = \sqrt{GM/R}$. In these equations R is the equatorial radius - rapidly rotating stars are not spheres and so one must be careful how one measures the “radius”. In the extreme limit of critical rotation, we expect a ratio of equatorial to polar radii of 1.5.

The above equation ignores other sources that can counter-balance gravity. For massive stars we’ve already encountered one important source: radiation pressure. Recall in Lecture 8 we defined $\Lambda = \frac{\kappa L}{4\pi GMc}$ as the Eddington factor. We can then define an “effective mass” that needs to be supported as $M_{\text{eff}} = M_*(1 - \Lambda)$, so that our definition of the critical velocity becomes:

$$V_{\text{crit}} = \sqrt{\frac{GM_*(1 - \Lambda)}{R_{\text{eq}}}} \quad (13.1)$$

Rotation induces efficient mixing via several processes, including shear caused by differential rotation and meridional circulation. The latter is driven by the local breakdown of radiative equilibrium. In the case of rigid rotation, the latter is referred to as Eddington-Sweet circulation, and the mixing timescale is proportional to the thermal (KH) timescale. Rotational mixing can bring nuclear processed material to the surface, and observations of processed material (in particular nitrogen) have been an important constraint on rotating stellar models.

Rotation has several important effects on the evolution of massive stars. Rotationally-induced mixing affects the abundance profile in the star, which affects the luminosity and to some extent the radius. Mixing with the core can bring fresh fuel to the center, prolonging the main sequence lifetime. Rotation generally results in a larger star overall, and hence lower density and lower T_{eff} compared to a non-rotating star. The details are complex and

depend on lots of uncertain physics. This general behavior is only a rough guide.

In normal stars, nuclear burning establishes a chemical stratification that gives rise to the usual core-envelope structure that then governs much of the subsequent evolution. However, when rotation rates are sufficiently high so that the timescale of rotationally induced mixing is shorter than the nuclear timescale on which the star usually establishes a chemical stratification, then the star evolves in a nearly chemically-homogeneous manner. What happens in this case is that the star at the end of core-H fusion finds itself on the helium main sequence. In other words, it does not evolve toward cooler temperatures in the usual way. This is a very distinct mode of evolution of massive stars, and is a fairly generic prediction of rapidly rotating stars. Stars in this phase have not yet been observed (the evolutionary track places them near the hydrogen main sequence so they are difficult to identify).

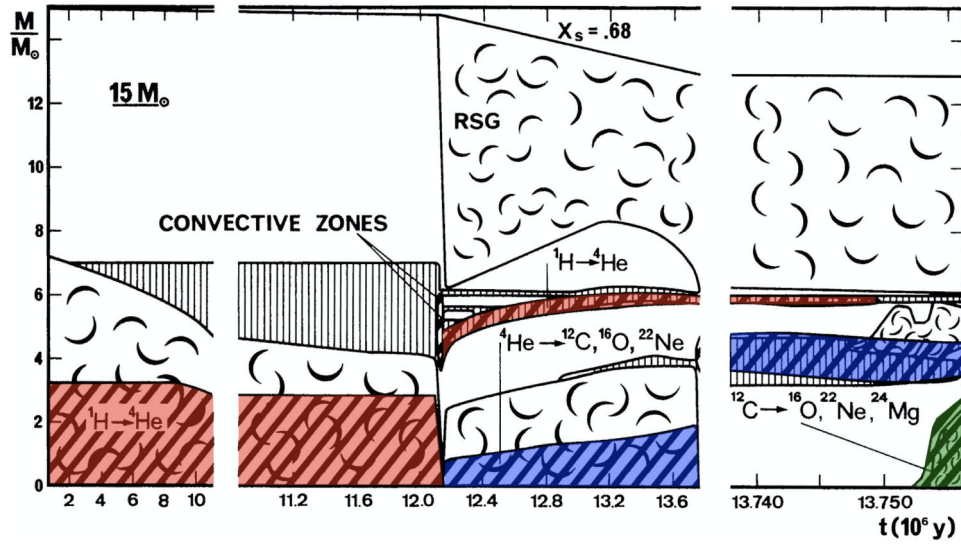


Figure 13.1: Kippenhahn diagram for a $15 M_{\odot}$ solar metallicity model. Curly regions indicate convection. Vertical hatched regions have abundances changed due to earlier periods of convection. Hydrogen, helium, and carbon burning are highlighted in red, blue, and green, respectively. Note the discontinuous x-axis. From Lamers & Levesque Fig 23.2.

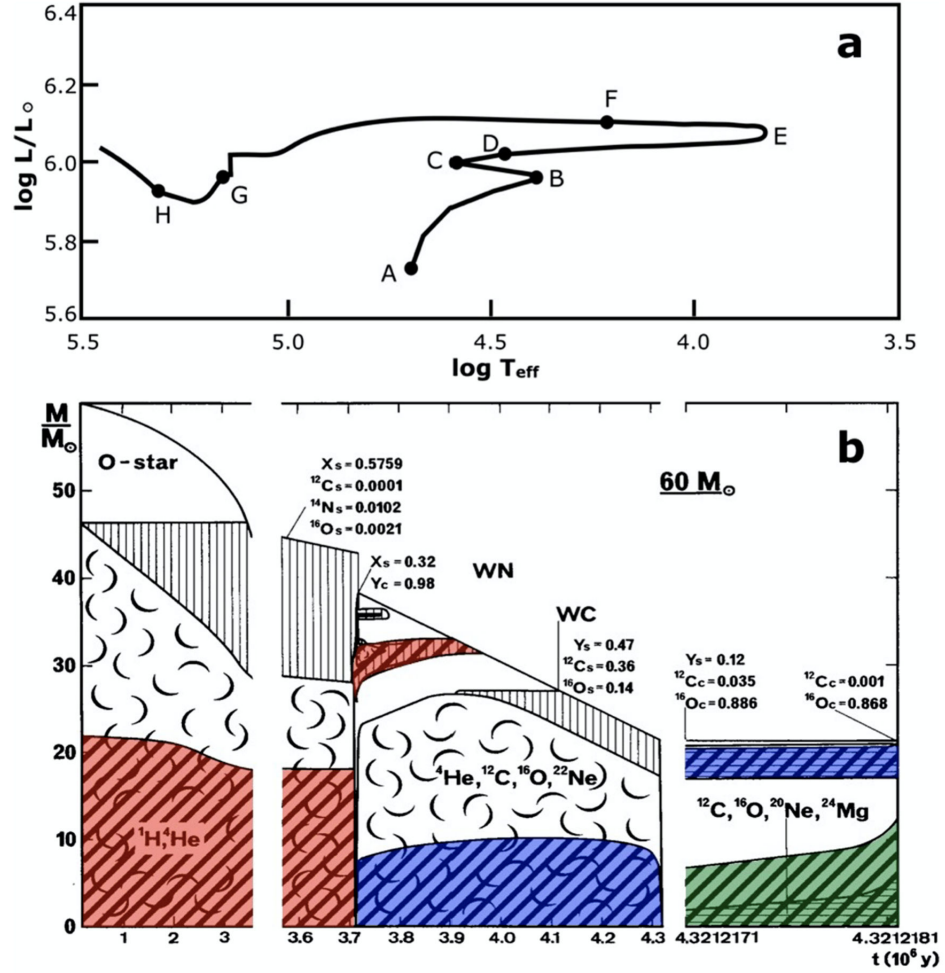


Figure 13.2: Evolution of a $60 M_{\odot}$ solar metallicity model with mass-loss and no rotation. Top panel shows the evolutionary track, bottom panel shows the Kippenhahn diagram, with labeling as in Figure 13.1. Notice the very strong mass-loss which reveals the interior of the star. From Lamers & Levesque Fig. 24.4.

Chapter 14

Binary Stars I

Binary star systems are common. The majority of massive stars ($> 10M_{\odot}$) are in binary systems, and amongst the low-mass stars, the binary fraction increases strongly with decreasing metallicity, potentially exceeding 50% at $[\text{Fe}/\text{H}] < -2$. Binarity can arise in one of two ways: a system can be born as a binary, or it can dynamically form as a binary. The latter channel is restricted to dense environments, e.g., in the central regions of globular clusters. The dynamical formation channel is relatively well-understood on theoretical grounds (and requires a third body to conserve energy), while the star-formation channel is relatively uncertain.

Binary stars play key roles in several astrophysical contexts. When observed as detached eclipsing binaries, they offer unique and powerful constraints on stellar parameters. The Hulse-Taylor pulsar is a binary neutron star system that provided the first convincing evidence of gravitational wave emission. Binary star mass-transfer gives rise to a wealth of observational phenomena, including low-mass and high-mass X-ray binaries and Algol systems. Binary star evolution is believed to play a critical role in creating the gravitational wave events now routinely detected by LIGO.

There are also a non-trivial fraction of triples and higher-order multiples. In most cases, a triple consists of a tight binary and a third star on a much longer period orbit, so in practice one can treat the tight binary as evolving on its own (this is known as a hierarchical triple). There are some situations in which triple-star evolution produces interesting phenomena.

Observationally there are several key distributions that characterize binary stars:

- The mass-ratio distribution. We define $q \equiv M_2/M_1$ where q is the mass ratio, M_1 is the primary and M_2 is the secondary mass in the system. The distribution is often written as a power-law, $f(q) \propto q^K$, and observationally $K \approx 0$, i.e., a flat distribution in q .

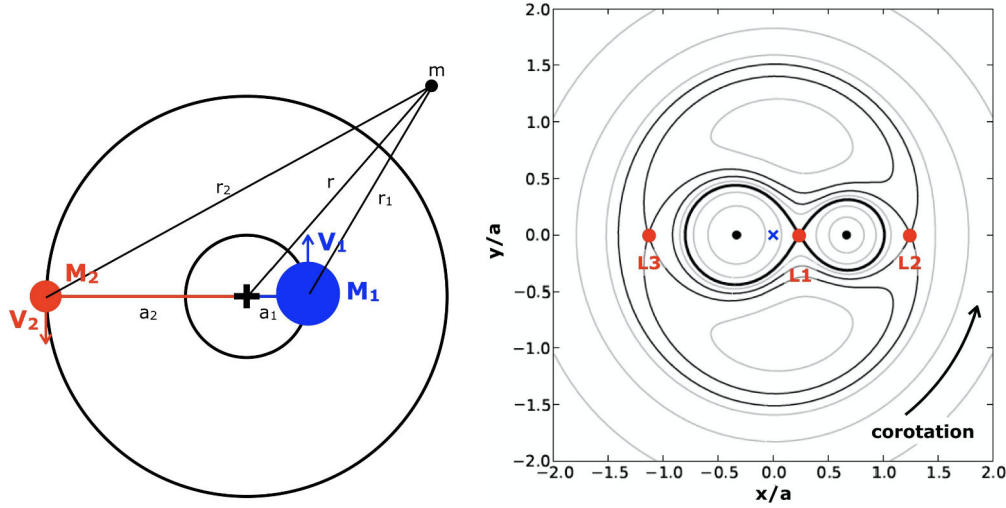


Figure 14.1: *Left Panel:* Geometry defining the key quantities for a binary star system orbiting a common center of mass (indicated by the “+” symbol). M_1 is the more massive star in the system. *Right Panel:* Iso-potential surfaces in the co-rotating frame. The blue cross is the COM, and the three red symbols denote three Lagrange points. Figures from Lamers & Levesque.

The details are of course more complicated; see Moe & Di Stefano (2017) for details.

- The period distribution, often expressed as $f(\log P) \propto (\log P)^\pi$. Empirically, $\pi \approx -0.5$ for $0.15 < \log P/\text{days} < 3.5$.
- The eccentricity (e) distribution. This is much less well-constrained observationally, and as a consequence one often assumes circular orbits.

Binarity will only affect the evolution of the stars if they can interact. We therefore need to understand the conditions under which stars may transfer mass between each other. This in turn requires some knowledge of the gravitational potential of a binary star system.

Let’s consider the properties of two stars in orbit around a common center of mass (COM; see Figure 14.1). Kepler’s third law tells us:

$$\left(\frac{2\pi}{P}\right)^2 = \omega^2 = \frac{G (M_1 + M_2)}{a^3} \quad (14.1)$$

where P is the period and ω is the angular velocity. The center of mass is defined by:

$$M_1 a_1 = M_2 a_2 = M_1 M_2 / (M_1 + M_2) a = \mu a \quad (14.2)$$

where $\mu \equiv M_1 M_2 / (M_1 + M_2)$ is the reduced mass and $a = a_1 + a_2$.

Consider potential surfaces in a *corotating* frame with angular velocity ω . The potential in this frame is:

$$\Phi(r) = -\frac{GM_1}{r_1} - \frac{GM_2}{r_2} - \frac{1}{2} \frac{G(M_1 + M_2)}{a^3} r^2 \quad (14.3)$$

where the last term represents the centrifugal acceleration and can be written as $\Phi_c = -\frac{1}{2}\omega^2 r^2$ (where r is the distance to the COM). This equation defines **equipotential surfaces**, which depend only on q (mass ratio) and P (period). Particles can freely move along equipotential surfaces (i.e., they move along orbit surfaces of constant potential).

Lagrange points are locations where $d\Phi/dr = 0$, i.e., where the force is zero. Such points are labeled as L_1 , etc (see Figure 14.1). L_2 is where many astronomical satellites are located, e.g., WMAP, Planck, Herschel, Gaia, JWST. The equipotential that intersects the L_1 Lagrange point is known as the **Roche lobe**. The Roche lobe is the tightest surface where particles can flow freely from one star to the other. A star dominates the gravitational potential within its Roche lobe. Material that is beyond L_2 and L_3 is bound to the binary system but not to either individual star. **Corotation** is the location where a free particle in orbit would have the same orbital frequency as the binary.

We can approximate the Roche radius (R_L) of star 1 as $R_1 \approx a(M_1/3M_2)^{1/3}$ which represents the balance between the tidal force and self-gravity. A more accurate expression for the Roche radius is provided in HW4.

Binarity will impact stellar evolution only if one or both stars fill their Roche lobe. At that point **mass transfer** will occur. We can classify binaries according to whether they are Roche lobe filling: 1) detached binary: neither star fills its Roche lobe; 2) semi-detached binary: one star is filling its Roche lobe; 3) contact binary: both stars are filling their Roche lobes. Note that these terms refer to the present state of the binary.

Mass transfer is conservative if $\dot{M}_1 = -\dot{M}_2$. In this case no mass is lost from the system and we have $\dot{J} = 0$ where J is the angular momentum of the system:

$$J^2 = M_1 a_1 v_1 + M_2 a_2 v_2 = Ga \frac{M_1^2 M_2^2}{M_1 + M_2} \quad (14.4)$$

(assuming circular orbits).

The assumption of conservative mass-transfer directly determines the evolution of the orbit, i.e., $a(t)$, $q(t)$. For example:

$$\frac{\dot{a}}{a} = 2 \frac{\dot{M}_1}{M_1} \left(\frac{M_1}{M_2} - 1 \right) \quad (14.5)$$

If M_1 is the donor, then when $\dot{M}_1 < 0$ (i.e., when M_1 is donating mass), the orbit will *shrink* if $M_1/M_2 > 1$. The minimum orbit separation is reached when $M_1 = M_2$. Subsequent mass transfer from M_1 will result in an increase in the orbital separation. This orbital evolution behavior is an important aspect of the evolution of interacting binaries.

Conservative mass transfer is a useful idealization, but it is unlikely to occur in most situations in nature. This is unfortunate, because non-conservative mass transfer is much more complicated to model. In the non-conservative case, some fraction of the mass lost by the donor is expelled from the system. This situation is more complicated because the details depend on how much angular momentum the expelled material carries away. The non-conservative case requires specification of at least two additional parameters: the fraction of mass expelled, and the specific angular momentum of the expelled material. Realistic simulations are required to determine these parameters.

There are three primary categories of mass-transfer binaries that depend on the evolutionary phase of the donor. These are known as Case A, B, C, and lead to very different evolutionary outcomes. The key issues to consider are how the stellar radius reacts to changes in mass, how the Roche lobe behaves as a function of mass, and the balance between the stellar and Roche radii. For the stellar radius response, we can expect very different behavior depending on the evolutionary state of the star.

Case 1: stable mass transfer: In this case the donor radius decreases as mass is lost, and it does so faster than the decrease in the Roche radius. The timescale for the star to inflate again and fill its Roche lobe is the nuclear timescale. Examples of stars in this configuration are binaries on the main sequence. Stars on the main sequence undergoing stable mass transfer are said to be undergoing **Case A** mass transfer.

Case 2: thermal-timescale mass transfer: In this case mass transfer is driven by the thermal readjustment of the donor, i.e., the mass-transfer rate saturates at $-M_d/\tau_{\text{KH},d}$ where M_d is the donor mass and $\tau_{\text{KH},d}$ is the thermal timescale of the donor. In this case the star is still in hydrostatic equilibrium. An example of this situation is **Case B** mass transfer in which a star has expanded after leaving the main sequence but has not yet reached the Hayashi line (e.g., stars in the Hertzsprung Gap in which H-shell fusion is occurring but the

star still has a radiative envelope).

Case 3: dynamically-unstable mass transfer: In this case the adiabatic response of the star to mass-loss is unable to keep the star within its own Roche lobe, leading to ever-increasing mass-transfer rates. This leads to runaway mass-loss on a dynamical timescale and probably leads to a common-envelope situation. We find this situation in stars that have deep convective envelopes, and in such cases we refer to **Case C** mass transfer. The reason that deep convective envelopes lead to runaway mass transfer can be understood as follows. Consider a star on the red giant branch. Such stars are on the Hayashi track (approximately constant T_{eff}) and we know from stellar evolution that luminosity is governed by their core mass: $L = f(M_c)$. *In this situation the radius is independent of the envelope mass.* Once mass-transfer begins, a decreases and so R_L decreases. But the stellar radius is approximately constant. In other words, the star does not respond to a loss of mass. The donor soon finds itself out of hydrostatic equilibrium and subsequent evolution occurs on a dynamical timescale. The process ends when the entire envelope is stripped, leaving behind a helium white dwarf in the case of low-mass stars or a Wolf-Rayet star for high mass stars.

Chapter 15

Binary Stars II

In the previous lecture we considered three types of mass-transfer binary systems: Case A, B, and C. Let's consider a few real-world examples of these cases.

Algol systems are binary systems in which a Roche-lobe filling giant is in orbit around a more massive main sequence companion. This configuration was puzzling to astronomers for a long time, as one expects the more massive star to evolve off the main sequence first. The puzzle was resolved when it was realized that Case A mass transfer could transfer substantial mass from a more massive star to a less massive companion, such that the initially less massive star becomes the more massive one, and the initially more massive star becomes less massive and evolves off the main sequence. This is a channel for creating **blue straggler stars** - stars that are hotter and more luminous than the main sequence turnoff point in star clusters.

Wolf-Rayet – O star binary systems form when an initially more massive star evolves across the Hertzsprung gap and fills its Roche lobe. Mass transfer occurs on a thermal timescale (Case B). Evolution is fast until $M_1/M_2 \approx 1$, then the binary expands and $\dot{M}$ decreases. The mass-loss strips the envelope of the more massive star, creating a Wolf-Rayet star (i.e., a star lacking a H envelope).

Common envelope systems form when mass transfer begins on the Hayashi line (Case C mass transfer). In this case runaway mass transfer occurs on a dynamical timescale. A situation arises in which the system literally has a common envelope and two distinct cores. There are only a few potential CE systems known. They should be a transient phenomenon, with timescales of weeks to months. Our best evidence for CE events is in the discovery of double compact objects with very tightly bound orbits.

Common envelope evolution is a fascinating process that is probably critical for the formation of at least some of the high-mass black holes now being detected by LIGO. The basic setup consists of two distinct cores orbiting within a low-density common envelope. The orbiting stars experience a drag force (friction) against the background envelope. The net effect is a transfer of orbital energy to the envelope. There are two possible outcomes: 1) the orbiting stars merge together before the envelope is ejected; 2) the envelope is ejected before the orbiting stars merge. Details depend on the energy in the envelope, E_{env} , and orbital energy, E_{orb} , as well as the uncertain physics of the energy exchange between the orbit and the envelope.

We can sketch the basic picture as follows. The overall efficiency of the CE process is:

$$\alpha_{\text{CE}} \equiv \frac{\Delta E_{\text{bind}}}{\Delta E_{\text{orb}}} \quad (15.1)$$

where ΔE_{orb} is the change in orbital energy at the beginning and end of the CE phase, and ΔE_{bind} is the binding energy of the envelope at the beginning of the CE phase. We have:

$$\Delta E_{\text{orb}} = \frac{GM_a M_{d,\text{core}}}{2a_f} - \frac{GM_a M_d}{2a_i} \quad (15.2)$$

where a_f and a_i are the final and initial orbital separations, M_a is the accretor mass, and M_d and $M_{d,\text{core}}$ are the original donor and donor core masses. Next we can estimate the binding energy of the donor envelope at the moment when it fits its Roche lobe (R_L):

$$\Delta E_{\text{bind}} = \frac{GM_d M_{d,\text{env}}}{\lambda R_{L,d}} \quad (15.3)$$

where λ is a dimensionless quantity encapsulating the core-envelope structure of the donor; this parameter can readily be computed from stellar structure models.

From these equations we can compute the ratio between final and initial separations:

$$\frac{a_f}{a_i} = \frac{M_{d,\text{core}}}{M_d} \frac{M_a}{M_a + 2M_{d,\text{env}}/(\alpha_{\text{CE}} \lambda r_{L,d})} \quad (15.4)$$

where $r_{L,d} = R_{L,d}/a_i$ is the relative Roche lobe size at the start of CE and is a function only of M_d/M_a . Usually the second term in the denominator is larger than the first, so we can approximate the above expression as:

$$\frac{a_f}{a_i} \approx \frac{1}{2} \alpha_{\text{CE}} \lambda r_{L,d} \frac{M_{d,\text{core}} M_a}{M_d M_{d,\text{env}}} \quad (15.5)$$

This is known as the “ $\alpha - \lambda$ ” formalism for CE evolution, where α_{CE} is the key free parameter. The state-of-the-art involves estimating this parameter based on 3D hydrodynamic simulations including radiative transfer, turbulence, and other potentially relevant processes.

There is no consensus yet on the outcome from these simulations, but values of $\alpha_{\text{CE}} \lesssim 1$ are generally found.

CE evolution is likely the cause of very short-period binaries containing one or more compact remnants. Examples include:

1. High-mass X-ray binaries (HMXB), in which a main sequence (e.g., O or B-type star) orbits a neutron star or black hole. The periods are very short, $\sim 1 - 2$ d. Mass-transfer results in accretion onto the compact object, which produces bright x-ray emission powered by accretion (see below). The most famous example of an HMXB is Cygnus X-1 (the first identified black hole candidate).
2. Low-mass X-ray binaries (LMXB), in which a low-mass star ($\lesssim 2M_{\odot}$) orbits a compact object (NS or BH). These also have short periods. Hundreds are known in the Galaxy. LMXBs have also been detected in other nearby galaxies.

There are many instances in astronomy in which the luminosity of an object is powered by **accretion**, including during the initial collapse of a protostar and accretion onto black holes. The basic idea is that energy is liberated as a mass m falls onto the surface of an object of radius R and mass M . The energy is usually released as radiation. We then have, for each mass m : $\Delta E_{\text{acc}} = GMm/R$. From this we can define an accretion luminosity:

$$L_{\text{acc}} = \frac{GM\dot{m}}{R} \quad (15.6)$$

We can also define a dimensionless efficiency:

$$\epsilon \equiv \frac{L_{\text{acc}}}{\dot{m}c^2} = \frac{GM}{Rc^2} \quad (15.7)$$

which can be very large. For a $1.4M_{\odot}$ neutron star with $R = 10^6$ cm, $\epsilon \approx 0.2$. Compare with the efficiency of fusion: $\epsilon = 0.007$. For black holes the “radius” is $R_S = 2GM/c^2$ which would suggest $\epsilon = 0.5$. However, the more relevant radius is the innermost stable circular orbit (ISCO), which, as the name suggests, is the smallest radius that can support a stable circular orbit. For a non-rotating BH this is $3R_S$, which would suggest a maximum efficiency of $\epsilon \approx 0.17$. It is an active area of research to determine the efficiency as a function of BH mass, spin, and properties of the accretion disk.

Recall the Eddington luminosity: $L_{\text{Edd}} = 4\pi cGM/\kappa$ where $\kappa = \sigma_T/m_H$. If we set $L_{\text{Edd}} = L_{\text{acc}}$ as the maximum possible accretion luminosity, then we can arrive at an expression for the

maximum possible accretion rate: $\dot{M}_{\text{max}} = 4\pi cR/\kappa$, which for Thomson opacity is $10^{-5}M_{\odot} \text{ yr}^{-1}$ for white dwarfs and $10^{-8}M_{\odot} \text{ yr}^{-1}$ for neutron stars. It is an active area of research to investigate accretion at or near L_{Edd} . In particular for accretion onto a black hole, one could imagine situations in which the material does not efficiently radiate energy, in which case the accretion rate could be much higher than L_{Edd} . Furthermore appealing to non-spherical geometries could allow photons to “escape” in a non-uniform manner (either perpendicular to a disk or via photon bubbles), thereby circumventing the Eddington limit.

There is a significant complication to the above story. Mass lost from the donor star typically has too much angular momentum to fall directly onto the surface of the accreting star owing to the angular momentum of the orbit. We can expect that the material will therefore form an **accretion disk**. The modeling of such phenomena is quite complicated and is an exciting area of current research.

Let’s consider a thin Keplerian disk in which the self-gravity of the disk and gas pressure within the disk are unimportant. Under these assumptions the rotational velocity is $v = \sqrt{GM/r}$ for a given radius r from the star of mass M . The specific angular momentum (the angular momentum per unit mass) is $j \equiv vr = \sqrt{GMr}$. Compare this to the maximum possible specific angular momentum at the stellar surface: $j = \sqrt{GMR}$. Since $r > R$, we have a problem: the angular momentum of the disk is generally much higher than the maximum possible angular momentum of the star. The total angular momentum of the system must be conserved, so how can we get material from the disk to the stellar surface?

To begin answering this question, consider two rings of material within the Keplerian disk. The inner ring is moving faster. Friction between the two rings will cause the inner ring to slow down and the outer ring to speed up. This results in angular momentum transfer from the inner to the outer ring. Because the specific angular momentum is $j = \sqrt{GMr}$, the loss of angular momentum will cause the inner ring to *sink inwards*. The net effect is that most angular momentum is transferred to the outer disk while the inner disk sinks to the center. But a key question remains: what physical mechanism is responsible for the transfer of angular momentum?

Here we must consider the possible transport processes available. A natural candidate is **viscosity**, ν (sometimes referred to as the momentum diffusivity). However, the molecular viscosity is far too small to act as the source of friction. Molecular viscosity is the result of molecular diffusion that transports momentum between layers of a flow; it therefor scales as $\nu \sim c_s \lambda$ where c_s is the thermal velocity (sound speed) and λ is the mean free path. The timescale associated with this molecular viscosity ($\tau \sim R^2/\nu$) is orders of magnitude too

long to drive the observed phenomena.

The gas in the disk is expected to be turbulent (i.e., to have large-scale motions not captured by rotation nor thermal motions). Turbulence can induce an effective “turbulent” viscosity that enables diffusion of angular momentum between two zones in the disk.

There is no general theory of turbulence, but there is a classic approach to this problem from Shakura & Sunyaev (1973). We let the viscosity ν scale as $\nu = \alpha c_s H$ where c_s is the sound speed, H is the scale height, and α is a free parameter. We shouldn’t expect turbulent velocities to be much higher than the sound speed, otherwise the turbulence would dissipate in shocks. And the largest turbulent eddy cannot be larger than the scale-height. This is known as the α -disk model for accretion disks. To give you a sense of the centrality of this model to accretion disk physics, the Shakura & Sunyaev paper has over 10,000 citations!

The most widely accepted explanation for angular momentum transport within disks is magnetic fields. Specifically, a process known as the magneto-rotational instability (MRI) proposed by Balbus & Hawley (1991). In this scenario, the disk is threaded with magnetic fields, which have a tension that is similar to a spring. The tension between two rings will cause the inner ring to lose angular momentum, slow down, and sink inward. This is a runaway process. Detailed simulations have been done to show that MRI can be an efficient transport process, and it has the significant advantage of being a physical model, unlike the α -disk.

Supplementary Material

Lets now explore the properties of accretion disks. Consider a disk in cylindrical geometry with radius R and surface mass density Σ . Now consider two thin annuli of matter on either side of the surface $R = \text{constant}$. We can express the viscous torque exerted on the inner annulus by the outer annulus as:

$$T(R) = 2\pi\Sigma R^2\nu R\frac{\partial\Omega}{\partial R} \quad (15.8)$$

where ν is the viscosity. Equation 15.8 can be derived from the Navier-Stokes equation, which describes the conservation of momentum of a fluid subject to viscosity. We can gain some insight into this equation by the following considerations. In the above equation $R\frac{\partial\Omega}{\partial R}$ is the velocity shear. The dynamic viscosity, μ , is defined as the ratio between force per unit area and the shear velocity, and the kinematic viscosity (ν above) is defined as $\nu \equiv \mu/\rho$ where ρ is the density. By combining these definitions we arrive at the equation for torque above. Recall that torque is the rate of change of the angular momentum, i.e., $T = dJ/dt$,

so if the torque is non-zero then there will be angular momentum flow within the disk.

Mass conservation within the disk implies:

$$R \frac{d\Sigma}{dt} + \frac{\partial}{\partial R}(R\Sigma u_R) = 0 \quad (15.9)$$

where u_R is the radial drift velocity.

Conservation of angular momentum implies:

$$R \frac{\partial}{\partial t}(\Sigma R^2 \Omega) + \frac{\partial}{\partial R}(\Sigma u_R R^3 \Omega) = \frac{1}{2\pi} \frac{\partial T}{\partial R} \quad (15.10)$$

If we assume that the disk mass is negligible compared to the central star, then the disk will exhibit Keplerian rotation ($\Omega = \sqrt{GM/R^3}$). Combining with our earlier expression for $T(R)$, we have the solution for a thin Keplerian disk:

$$\frac{\partial \Sigma}{\partial t} = \frac{3}{R} \frac{\partial}{\partial R} \left[R^{1/2} \frac{\partial}{\partial R} (\nu \Sigma R^{1/2}) \right] \quad (15.11)$$

$$u_R(R, t) = -\frac{3}{\Sigma R^{1/2}} \frac{\partial}{\partial R} (\nu \Sigma R^{1/2}) \quad (15.12)$$

We can understand what this solution means by assuming $\nu = \text{constant}$ and an initially infinitesimally thin ring with mass m and radius R_0 :

$$\Sigma(R, t=0) = \frac{m \delta(R - R_0)}{2\pi R_0} \quad (15.13)$$

where $\delta(y)$ is the δ function. Define a dimensionless radius $x \equiv R/R_0$ and a dimensionless time $\tau \equiv t(12\nu/R_0^2)$. Then the solution of the above equations for the surface density evolution is:

$$\Sigma(x, \tau) = \frac{m}{\pi R_0^2} \tau^{-1} x^{-1/4} \exp \left[-\frac{1+x^2}{\tau} \right] I_{1/4} \left(\frac{2x}{\tau} \right) \quad (15.14)$$

where $I_{1/4}$ is a modified Bessel function. The solution is shown in the slides. At early times the surface density is a narrow Gaussian around $R = R_0$. At later times most of the mass has lost angular momentum and moved inward, but there is a tail of material that moves out to larger and larger radii which carries away most of the angular momentum of the system.

The next step is to let $D(R)$ be the rate per unit time per unit area at which the kinetic energy of rotation is dissipated into heat by viscosity:

$$D(R) = \frac{1}{2} \nu \Sigma \left(R \frac{\partial \Omega}{\partial R} \right)^2 \quad (15.15)$$

This is known as the viscous dissipation rate. We can then compute the total disk luminosity as:

$$L_{\text{disk}} = 2\pi \int_0^\infty D(R) R dR \quad (15.16)$$

which for a steady-state Keplerian disk is just:

$$L_{\text{disk}} = \frac{1}{2} \frac{GM\dot{m}}{R} \quad (15.17)$$

Notice that L_{disk} is exactly 1/2 of the total accretion luminosity. Evidently half of the accretion luminosity is emitted during the passage of the gas through the accretion disk; the other half must be emitted when the gas makes the transition from the inner edge of the disk to the surface of the compact object (unless the compact object is a BH, in which case at least some of the energy could go directly into the BH without being radiated).

Chapter 16

Stellar Variability

All stars are variable stars. In some cases the variability amplitude is parts per million (ppm), as in the Sun, while in other cases the amplitude of variability can exceed several magnitudes. Some variability is regular (periodic), with timescales from hours to years, while other variability is irregular or eruptive. Variation can be extrinsic (e.g., eclipsing binaries, microlensing), or intrinsic. As we will see, there are many physical mechanisms that give rise to stellar variability. Variability offers fundamental insight into the stellar interior. Beyond their importance in stellar astrophysics, some variables, such as Cepheids and RR Lyrae, are standard candles and have therefore played an essential role in mapping the extent of the Galaxy and the scale and expansion of the Universe.

The study of variable stars has a long and rich history. Several stars are visible as naked-eye variables, including Mira and Algol. It is unclear to what extent this was appreciated in ancient civilization, as there is little or no direct reference to their variability in the historical record. The first recorded observations of variability appear around 1600 (prior to the invention of the telescope). With even crude telescopes, many variable stars were identified in the 1700s (and many more in subsequent centuries). Initially, it was unclear if all variable stars were eclipsing binaries or not, as this was the most obvious explanation for variability. Observed variation in color with the phase of variability was an important clue, as well as the realization that many variables are *giants* and so could not have a companion at the necessary separation (via Kepler's law). This conclusion was reached by Harlow Shapley, Director of the Harvard College Observatory from 1921–1952. The Observatory has a long history in studying variable stars, including Henrietta Swan Leavitt, who famously discovered the period-luminosity relation for Cepheid variables (now known as the Leavitt Law).

There is a wide range of intrinsic variability mechanisms:

- Surface magnetic fields and spots (BY Draconis variables). Star spots combined with stellar rotation produce periodic variability.
- Flares due to magnetic fields (UV Ceti variables) produce irregular variability.
- Pre-main sequence stars undergoing rapid accretion (FU Ori [FU Orionis] and EXor variables), and T Tauri stars.
- Massive evolved stars: Wolf-Rayet (WR) and luminous blue variable (LBV) stars. The underlying physical mechanism is unknown, but is perhaps related to super-Eddington luminosities or other instabilities in the envelope.
- Mass-loss leading to rapid formation of dusty carbon, which obscures the optical light. Leads to rapid dimming for timescales of months to years. These are known as R Coronae Borealis (R Cor Bor) stars.
- Explosions and eruptions: novae, supernovae, cataclysmic variables.
- Regular pulsators. This is a large category and contains two basic types: radial and non-radial pulsators. Examples include Cepheids, RR Lyrae, Mira, δ Scuti, β Cephei, α Cygni, and ZZ Ceti.

See Figure 16.1 for the locations of various common variable stars in the HR diagram. Notice the existence of an **instability strip** within which many classes of variable stars are found.

Before considering *why* a star might pulsate, let's explore the pulsation period of a star and what it might tell us about the stellar interior.

Consider the sound speed, which is defined as $c_s^2 \equiv (\partial P / \partial \rho)_s$ where the derivatives are taken at constant entropy because a sound wave is usually sufficiently fast so that its propagation can be considered as an adiabatic process. For adiabatic processes we have $PV^\gamma = \text{constant}$, where γ is the adiabatic index, which can be expressed as $\gamma = (\partial \ln P / \partial \ln \rho)_s$. We can combine these various factoids of thermodynamics to write the speed of sound as $c_s^2 = \gamma P / \rho$ where $\gamma = 5/3$ for a monatomic ideal gas. We can therefore express the sound speed as $c_s^2 = \frac{5kT}{3\mu m_H}$.

We can then imagine that the period (Π) of variability is related to the time it takes a sound wave to cross a stellar diameter: $\Pi = 2R/c_s$.

For illustrative purposes let's consider the star δ Cephei, the prototype Cepheid variable. We know from interferometry that its angular size is 1.5 mas and its parallax is $\pi = 3.66$

mas, so its radius is $R = 45R_\odot$. Now what temperature should we choose for computing c_s ? Lets guess the temperature of the 2nd ionization of He (He II), i.e., $\approx 45,000$ K (we'll see later why this temperature is relevant). This implies $c_s \approx 32 \text{ km s}^{-1}$ and hence $\Pi = 22$ days. The observed period is $\Pi_{\text{obs}} = 5.4$ days. This is within an order of magnitude, but we can do better!

Rather than trying to adopt a single temperature, let's remember that the wave is traveling through the star, and will therefore be experiencing a range of temperatures, densities, and pressures. The period should instead be computed as an integral through the star:

$$\Pi = 2 \int_0^R \frac{dr}{c_s(r)} = 2 \int_0^R \frac{dr}{\sqrt{\gamma P(r)/\rho(r)}} \quad (16.1)$$

We can make progress by using the results of realistic stellar models (e.g., **MESA**), simple models (e.g., polytropes), or *really* simple models, e.g., a homogeneous model in which $\rho(r) = \rho_0$. For the homogeneous model, hydrostatic equilibrium implies:

$$\frac{dP}{dr} = -\frac{Gm\rho_0}{r^2} = -\frac{4\pi\rho_0^2 Gr}{3} \quad (16.2)$$

where for the second equality we set $m = 4/3\pi r^3 \rho_0$. Integrating and enforcing the boundary condition $P(R) = 0$ at the surface R leads to:

$$P(r) = \frac{2\pi\rho_0^2 G}{3} (R^2 - r^2) \quad (16.3)$$

Therefore:

$$\Pi = 2 \left(\frac{3}{2\pi^5/3\rho_0 G} \right)^{1/2} \int_0^R \frac{dr}{\sqrt{R^2 - r^2}} \quad (16.4)$$

the integral above is equal to $\pi/2$, so we arrive at our final expression:

$$\Pi = \sqrt{\frac{9\pi}{10G\rho_0}} \propto \rho_0^{-1/2} \quad (16.5)$$

We return to δ Cephei, which has a mass of $5M_\odot$, and therefore a mean density of $\rho_0 \approx 10^{-4} \text{ g cm}^{-3}$. The inferred sound crossing time in this case is $\Pi = 8$ days, closer to the observed value (5.4 days). Much better agreement can be obtained with a realistic stellar model.

For Mira, an AGB star which has $M = 1.2M_\odot$ and $R = 370R_\odot$, the predicted period is 412 days, while the observed pulsation period is 332 days. This is reasonably close for such a simple model.

Summary: the pulsation period probes the integral of the sound speed profile within the star and is hence a unique observational probe of the stellar interior.

Note that our derivation of the characteristic pulsation period looks very similar to the dynamical time: $t_{\text{dyn}} \propto \sqrt{1/G\rho}$. Recall that in most situations in stellar evolution we have the following temporal hierarchy: $t_{\text{dyn}} \ll t_{\text{KH}} \ll t_{\text{nuc}}$. This means that we can describe oscillations and pulsations with a static (non-evolving) stellar model. This is an important simplification. The oscillation code **GYRE**, for example, takes snapshot outputs from **MESA** for analysis. It does not need to consider the time evolution of the star as it computes pulsation frequencies.

We have only discussed pulsation periods (or equivalently, frequencies) – what about the amplitudes? It turns out that those are *much* more difficult to predict and are not amenable to neat back-of-the-envelope calculations. In fact, in many analyses the amplitudes are either ignored altogether or marginalized out of the results.

Stars can vibrate in many modes, much like vibrations on a string (or, if you want to 2D analogy, a drumhead). The period we have been focusing on thus far is known as the *fundamental mode*. There exists a hierarchy of modes: the first overtone (in which the period is exactly one half of the fundamental mode), second overtone, etc. Stars can (and sometimes do) pulsate in one or more overtone. Sometimes multiple modes are excited at the same time.

So far we have been discussing on what timescales a star might oscillate. Now let's turn to *why* a star would oscillate. This requires introducing the concept of **driving mechanisms**. To understand the concept, consider the air in your room. There is an associated sound speed, but no waves will travel unless there is some disturbance (e.g., your voice), which we call a driving mechanism. The same basic principle holds in stars (and all other astrophysical environments for that matter). In most situations we require the driving mechanism to be continuous in time, e.g., to continuously deposit energy into the pulsation. Any real world process is dissipative (e.g., by radiative diffusion or viscosity), which will damp the perturbation if not replenished. Consider wind blowing on the surface of a lake. The wind is a continuous driving mechanism that excites waves on the surface.

There are two principle ways to perturb in a star from its equilibrium solution: 1) modulate the energy production; 2) modulate the transport of radiation.

The first mechanism is known as the ϵ **mechanism**. It is always a source of oscillatory behavior because a compression increases the temperature, which increases ϵ (the energy production rate). This mechanism tends to be very weak in practice, but may be relevant in very massive stars.

The second mechanism is known as the κ **mechanism** (for opacity). This is *not* always a source of pulsation; it depends on the density and temperature of the layer. It turns out that this mechanism is the underlying cause of many radial pulsators including RR Lyrae, Cepheids, and Miras. The driving layer is associated with zones in the star where either H I or He II are partially ionized (which occurs at $T \approx 10,000$ K and $45,000$ K, respectively). The He II partial ionization zone is responsible for δ Scuti, RR Lyrae, ZZ Ceti, and Cepheid variables and defines an **instability strip** in the HR diagram.

Let's discuss how the κ mechanism works. In a fully ionized or neutral layer of a star, if a layer is adiabatically compressed, then the temperature will increase as $T \sim \rho^{2/3}$. Kramers opacity relation ($\kappa \sim \rho T^{-7/2}$) then tells us that the opacity will decrease because as $\kappa \sim \rho^{-4/3}$. The net effect is that when a layer is compressed radiation more easily escapes the layer, and the compression will damp out.

The situation in a partially-ionized zone can be very different. In this case, as a layer is compressed, part of the energy goes into ionizing the gas instead of raising the temperature. If we express the relation between the density and temperature in terms of the adiabatic index γ then in general we have $T \sim \rho^{\gamma-1}$, so Kramers relation can be expressed as $\kappa \sim \rho^{(9-7\gamma)/2}$, and therefore κ will increase during a compression if $\gamma < 9/7$. If the opacity increases during compression, then radiation flow is trapped in the compression layer. The net effect is that energy is input into the layer during a compression. In other words, we have found a driving mechanism that is able to continually inject energy into the pulsation. The compression will have more thermal energy than its surroundings, and will therefore expand, repeating the cycle. In partially-ionized zones the adiabatic index is in the range $1.19 < \gamma < 1.29$ (note that $9/7 \approx 1.29$) and so these zones are capable of driving oscillations.

Notice that what we really require is for the opacity to increase during a compression. This can occur outside of partial-ionization zones in layers where the opacity is not well-described by Kramers relation and is instead an increasing function of temperature. This occurs at the iron opacity peak at $\approx 1.8 \times 10^5$ K (caused by M-shell transitions of iron ions), and for the H^- opacity at low temperature. The latter is believed to be responsible for Mira pulsations. We can write the condition more generally as:

$$\left(\frac{d\ln\kappa}{d\ln P} \right)_{\text{ad}} = \left(\frac{\partial \ln\kappa}{\partial \ln P} \right)_T + \left(\frac{\partial \ln\kappa}{\partial \ln T} \right)_P \left(\frac{\partial \ln T}{\partial \ln P} \right)_{\text{ad}} \equiv \kappa_P + \kappa_T \nabla_{\text{ad}} > 0 \quad (16.6)$$

where κ_P and κ_T are shorthand for the partial derivative of opacity with respect to $\ln P$ and $\ln T$, respectively. For an ideal gas with Kramers opacity we have $\nabla_{\text{ad}} = 0.4$, $\kappa_P = 1$ and $\kappa_T = -4.5$ and hence $\kappa_P + \kappa_T \nabla_{\text{ad}} = -0.8 < 0$ and hence the star will not pulsate. In the partial ionization zones ∇_{ad} can be as low as 0.1 and hence the above expression can be > 0

for Kramers opacity, enabling pulsation. For H^- opacity we have $\kappa_T > 0$, which will also result in the above expression being > 0 .

What determines the blue and red edges of the instability strip? Lets consider the thermal timescale *above* the ionization zone, $t_{th,i}$. If $t_{th,i} \ll \Pi$, then the radiative gradient will quickly compensate for any change in opacity. In other words, if the thermal time is much shorter than the pulsation time, then the star can readjust before the pulsation has a chance to affect the emergent luminosity. This occurs when the partial ionization zone is near the surface, where the thermal time above the ionization zone is small. The partial ionization zone moves outward toward the surface for hotter stars. This sets the blue edge of the instability strip. What sets the red edge is a bit more ambiguous, but basically if $t_{th,i} \gg \Pi$ then the ionization zone is too deep to affect the luminosity. The instability strip can be understood as the region in the HR diagram where the thermal time above the partial ionization zone is comparable to the star's pulsation period, i.e., $t_{th,i} \approx \Pi$.

From Figure 16.1 it is obvious that many variable classes occur within the instability strip at distinct evolutionary phases. δ Scuti variables are A type main sequence stars. RR Lyrae are low mass core helium burning stars (i.e., horizontal branch stars) that intersect the instability strip, while (Type I) Cepheids are higher mass stars in the blue loop phase. RV Tauri variables (not shown in Figure 16.1) are post-AGB stars crossing the instability strip.

Let's return to Cepheids. There are in fact two types of Cepheids, Type I (Classical Cepheids or δ Cepheids) and Type II (W Virginis Cepheids). Type I periods are intermediate mass stars in the blue loop phase of their evolution. Type II stars are lower mass stars that have completed core helium burning on the horizontal branch and are now ascending the AGB. At a fixed period the Type I variables are more luminous than the Type II Cepheids. Confusion between these two types caused Hubble to severely over-estimate the Hubble constant.

Each type of Cepheid obeys a period-luminosity relation, and the relation for the Type I Cepheids is now known as the Leavitt Law. A great deal of effort has been expended to empirically determine this relation. The challenge is three-fold: observations spanning days to months are required in order to sample the range of periods expected. Second, distances to the reference sample must be known in order to determine luminosities. This is now routinely done with parallaxes for stars in the Galaxy. Third, extinction is an important source of systematic uncertainty. Recent work has focused on determining the Leavitt Law in near-IR passbands to minimize the effect of dust. The relation was originally measured in the Small and Large Magellanic Clouds in 1912 by Henrietta Swan Leavitt, while an

astronomer at the Harvard College Observatory.

Cepheids reside within the instability strip, and as such the physical mechanism for their pulsation is the κ mechanism driven by the partial ionization of He II. The period-luminosity relation is ultimately a mass sequence. Cepheids obey a mass-luminosity relation (due to standard stellar evolution arguments), and as we saw earlier the pulsation period is proportional to $\rho^{-1/2}$. More massive Cepheids are more luminous, less dense, and have longer periods. Obtaining a detailed, quantitative match to the observations remains challenging.

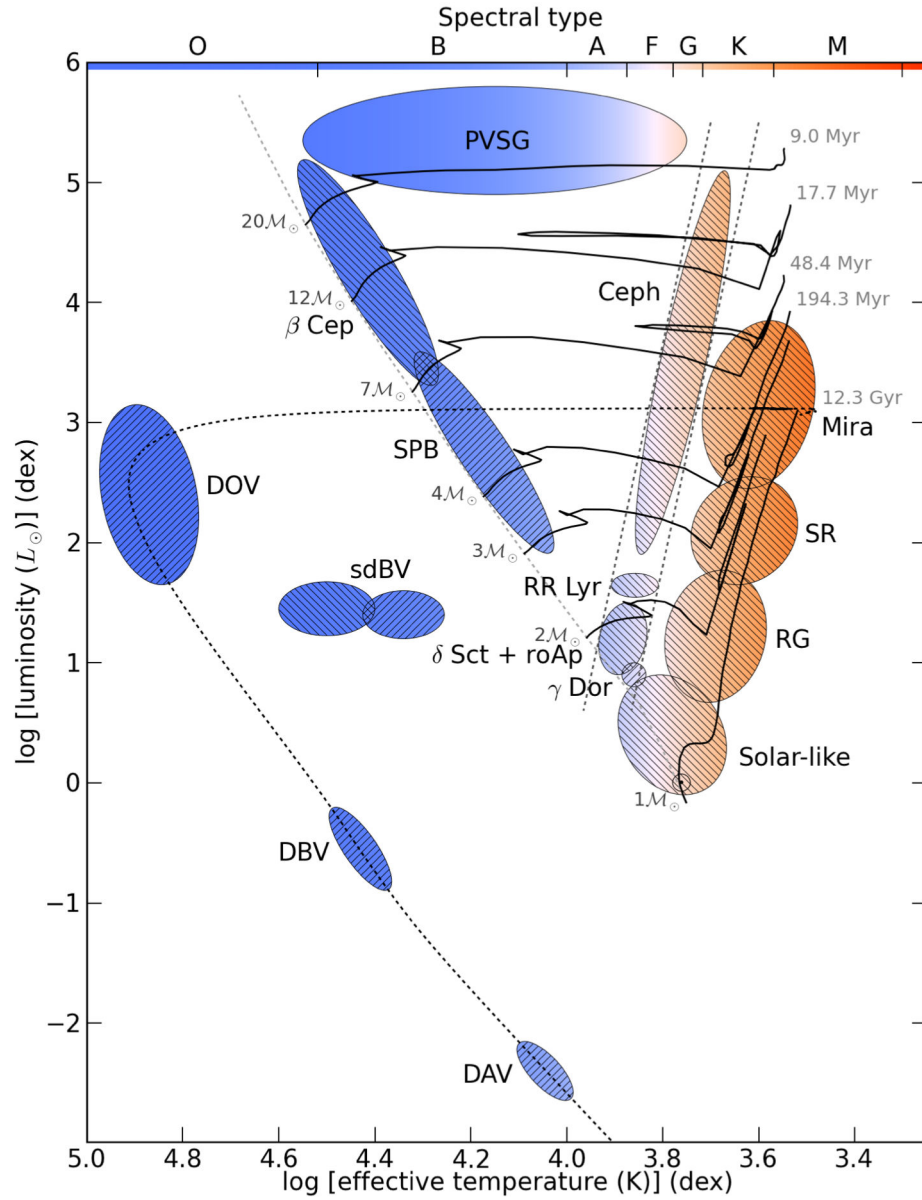


Figure 16.1: Location of common variable star types in the HR diagram. Hatched marks moving down to the right indicate stars where pressure modes are the dominant oscillation type, while hatched marks moving down to the left indicate gravity modes as the dominant type. From Aerts (2019).

Chapter 17

Asteroseismology

Asteroseismology is the study of stellar oscillations. Until relatively recently, such oscillations were only detectable in the Sun, and in this context is referred to as helioseismology. The first detected oscillations in the 1950s had a period of 5 min and a velocity amplitude of ≈ 500 m/s. This behavior emerges from the combination of $\sim 10^7$ modes¹ each with amplitudes of ~ 10 cm/s in velocity and ~ 1 ppm in flux (“ppm”, or parts per million, describes the fractional change in flux). In general the oscillations will induce changes in both brightness and velocity at the surface.

There are two significant hurdles to overcome in observing asteroseismology in stars. First, the amplitudes of flux variations are in the $\sim 1 - 100$ ppm range, so very stable, precise photometry is required. Owing to changes in observing conditions from the ground, space-based observations are a necessity in most cases. Second, the aliasing² induced by non-continuous observing of sources from the ground results in serious challenges in interpreting the mode spectrum. Early space missions from the Canadians and the French (MOST in 2003 and CoRoT in 2006, respectively) ushered in the beginning of the space-based asteroseismology revolution. The true revolution came with NASA’s Kepler mission, launched in 2009. Known primarily for its exoplanet science, Kepler’s photometric precision enabled a revolution in asteroseismology as well (100 times more precise than CoRoT). TESS (launched by NASA in 2018) is continuing the legacy, though in a somewhat different direction (shorter cadence,

¹A “mode” refers to a particular oscillation that can be characterized by a frequency, amplitude, phase, degree n , and, in the case of non-radial models, the spherical harmonics of degree l and azimuthal order m .

²Aliasing is the general phenomenon in which undersampling of a continuous datastream results in different signals being indistinguishable. For example, if a star was observed for one eight hour night, it would be impossible to distinguish periodicity on timescales longer than 8 hr. Furthermore, if the data were collected in 1 min exposures, it would be impossible to detect periodicity on < 1 min timescales.

brighter targets, but all-sky compared to Kepler).

There are two types of waves that we will consider in the context of stellar oscillations, both of which are standing waves. **Pressure waves**, i.e., acoustic or compression waves, are waves where the restoring force is provided by pressure. In the case of **gravity waves**, buoyancy provides the restoring force. An example of gravity waves are ocean waves. Gravity waves are *not* the same as gravitational waves. We'll discuss the latter in a later lecture.

Recall from our discussion of convection that we arrived at a condition for convective stability that appealed to the Brunt-Väisälä frequency:

$$N^2 \equiv g \left[\frac{1}{\gamma} \frac{d \ln P}{dr} - \frac{d \ln \rho}{dr} \right] \quad (17.1)$$

where $N^2 < 0$ implies that the region is unstable and convection will occur, while $N^2 > 0$ implied that the region is stable. It is the latter case that gives rise to gravity waves. In other words, a fluid element that is displaced will oscillate about that displacement. Gravity waves can therefore only exist in radiative regions (where the region is stable to such displacements). In the Sun, gravity waves are confined to the core. Gravity waves may play an important role in massive star evolution, where the amplitude of such waves is large and dissipation of the wave energy is possible.

A key feature of pressure waves is that non-radial modes will be confined to a certain zone of the interior depending on the angular degree, l of the mode. This is a consequence of refraction due to radial variation in the sound speed, c_s (see Figure 17.1 below). As a consequence, the behavior of the angular orders is a powerful probe of the stellar interior. In practice, high-order angular modes (higher l) are prone to cancellation effects in disk-averaged observations of stars: at higher l there are more independent modes across the stellar surface that are averaged away when viewing a star in total. This is not the case for the Sun, where we can observe the pulsations across the surface. The result is that we typically only observe the first few angular orders in stars beyond the Sun.

From the time-dependence of the flux one can perform a Fourier Transform to obtain the power spectrum of brightness variations. An example is shown in Figure 17.2. Predictions for the full spectrum of oscillations requires a realistic input model (e.g., from **MESA**) and full linear perturbation theory. In practice, this means employing the full time-dependent equations of mass, momentum, and energy for stellar structure.

Comparison between models and data is very cumbersome (as one can infer from the richness and complexity of the data in Figure 17.2), and is often referred to as “boutique” analysis, because the comparisons are done by hand, as opposed to an automated fashion. As a

consequence, the bulk analysis of modern data sets relies on simpler quantities that can be measured automatically. These are known as the **asteroseismic scaling relations**, and have become the standard way of analyzing data from the two large asteroseismic datasets from the Kepler and TESS satellites.

The acoustic cutoff frequency is defined as

$$\nu_{\text{ac}}^2 \equiv \frac{c_s^2}{16\pi H_P^2} \left(1 - 2 \frac{dH_P}{dr} \right) \quad (17.2)$$

where H_P is the pressure scale height. This is similar to the fundamental pulsation mode we discussed in the last lecture, except that we replace the stellar radius R with H_P . The reason for this has to do with the excitation mechanism: for solar-like oscillations the excitation is convection, and the modes most likely to be excited are those near the convective turnover timescale, which is in turn related to the pressure scale height. Oscillations with wavelengths much smaller than H_P do not reflect at the surface and therefore travel freely into the stellar atmosphere, resulting in strong damping.

propagate.

Empirically, one observes that ν_{ac} is similar to the frequency of maximum oscillation power, ν_{max} (see Figure 17.2 for an example). We therefore have:

$$\nu_{\text{max}} \sim \nu_{\text{ac}} \propto \frac{c_s}{H_P} \propto \frac{g}{\sqrt{T_{\text{eff}}}} \propto \frac{M}{R^2 \sqrt{T_{\text{eff}}}} \quad (17.3)$$

We then adopt the Sun as the reference point, so that:

$$\frac{\nu_{\text{max}}}{\nu_{\text{max},\odot}} = \left(\frac{M}{M_{\odot}} \right) \left(\frac{R_{\odot}}{R} \right)^2 \left(\frac{T_{\text{eff},\odot}}{T_{\text{eff}}} \right)^{1/2} \quad (17.4)$$

The large frequency separation, $\Delta\nu$, is the separation between modes of the same degree and consecutive radial order. $\Delta\nu$ is therefore inversely proportional to the acoustic radius. For vibrations on a string, the period separation between consecutive radial orders is proportional to $\int dr/c_s$, so:

$$\Delta\nu = \nu_{n-1,l} - \nu_{n,l} = \left[\int_0^R \frac{dr}{c_s} \right]^{-1} \propto \sqrt{\bar{\rho}} \propto \sqrt{\frac{M}{R^3}} \quad (17.5)$$

where the constants are again determined by scaling to the Sun:

$$\frac{\Delta\nu}{\Delta\nu_{\odot}} = \left(\frac{M}{M_{\odot}} \right)^{1/2} \left(\frac{R_{\odot}}{R} \right)^{3/2} \quad (17.6)$$

We now have two observables, $\Delta\nu$ and ν_{max} , and two unknowns, M and R . Asteroseismology therefore enables a measurement of the mass and radius of a star. Recall that these two

quantities are very difficult to determine by any other means, except in the case of detached eclipsing binaries. Quoted uncertainties of M and R measured in this way are frequently $< 10\%$. As you can probably imagine, the issues with this approach are the systematic uncertainties associated with the relatively simplistic sketch given above, and the assumption that the proportionality constant is indeed a constant that can be determined by comparison to the Sun. In reality, comparison to detailed “boutique” modeling of stars with high-quality data have revealed significant deviations from these scaling relations, and much work has gone into correcting them.

Rotation will, via the Coriolis force split modes (owing to the additional degree of freedom provided by the azimuthal order m), which allows a probe of the rotation in the core of the star. This requires core modes to escape. While core modes are not observed in the Sun, they have been observed in RGB stars.

What mechanism is exciting the modes? As we saw in the previous lecture, some mechanism must exist to drive the perturbations. The κ mechanism (driven by either local peaks in the opacity curve or partial ionization zones) is a means for exciting modes that will be long-lived. These are called undamped oscillations. In contrast, turbulent motions of convective regions can excite pressure waves. In this case each mode has a relatively short lifetime (days to months), and hence their mode cycle is regularly interrupted. The challenge in this case is that when the mode is re-excited it will have a different (random) phase, and this can significantly complicate the interpretation of the data. It is generally accepted that solar-like oscillations are stochastically excited by convection in the outer layers.

We are now going to explore the behavior of oscillation modes in the interior of stars in greater detail. In the supplementary material below we derive the equation of motion for linear adiabatic stellar pulsations:

$$\frac{d^2 \xi_r}{dr^2} = -\frac{\omega^2}{c_s^2} \left(\frac{N^2}{\omega^2} - 1 \right) \left(\frac{S_l^2}{\omega^2} - 1 \right) \xi_r. \quad (17.7)$$

where ξ_r is the radial displacement, N is the Brunt-Väisälä frequency, ω is the frequency of the oscillation, and S_l is the characteristic acoustic frequency (also known as the Lamb frequency): $S_l^2 \equiv l(l+1)c_s^2/r^2$. The solutions are oscillatory when the right hand side is negative and exponential when the right hand side is positive. Oscillations will therefore occur when two conditions are met: i) $|\omega| > |N|$ and $|\omega| > S_l$, or ii) $|\omega| < |N|$ and $|\omega| < S_l$. Within these regions the oscillations are trapped and the boundaries of the trapped region are defined by where the RHS is zero. Outside of these conditions the perturbations are

exponential, and in the case where they exponentially decay (which is the usual case) they are called evanescent modes.

Figure 17.3 is an example of a propagation diagram. The Brunt-Väisälä frequency, N_{BV} , is plotted as a function of stellar radius, and is compared to the Lamb frequency, S_l for select l -orders. Frequencies lower than N_{BV} are stable g -mode oscillations. These are confined to the core for the Sun because the outer layers are convective, where a g -mode is unstable (evanescent). Pressure wave (p -mode) oscillations are confined between the surface and a turnaround radius, r_t . This region is called the resonant cavity for the wave, and the turnaround radius is where the RHS of Eqn 17.7 is zero, which can be determined by setting $S_l = \omega$, which leads to:

$$\frac{c^2(r_t)}{r_t^2} = \frac{\omega^2}{l(l+1)} \quad (17.8)$$

This reason for this confinement is the refraction of sound waves discussed earlier. For example, the dashed line in Figure 17.3 shows that a p -mode at $l = 20$ and $\nu = 2000\mu$ Hz is confined to a turnaround radius of $r/R > 0.6$.

The key features of diagrams like Figure 17.3 is that g -modes are confined to the radiative layers and generally have lower frequencies than the p -modes; the latter span a much wider range of frequencies. The maximum depth probed by a p -mode depends on the mode degree l . The detailed behavior depends strongly on the interior structure of the star. This is advantageous, because it means that observations of the mode structure can provide sensitive constraints on the interior structure of stars.

Supplementary Material:

The following notes provides more a more detailed discussion of how oscillation modes are derived for stars.

Throughout the course we have dropped the time-dependence of our governing equations. In order to study time-dependent oscillations, we must return to the full equations of hydrodynamics expressing conservation of mass, momentum, and energy. Specifically, conservation of mass implies:

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{v}) = 0 \quad (17.9)$$

where $\rho(\mathbf{r}, t)$ and $\mathbf{v}(\mathbf{r}, t)$ are the local density and velocity at position vector $\mathbf{r}$ and time t . The equation of momentum conservation is:

$$\rho \frac{\partial \mathbf{v}}{\partial t} + \rho \mathbf{v} \cdot \nabla \mathbf{v} = -\nabla P - \rho \nabla \Phi + \rho \mathbf{f} \quad (17.10)$$

where Φ is the gravitational potential (satisfying Poisson's equation), P is the pressure, and $\mathbf{f}$ is the body force per unit mass. In the above equation we have ignored viscosity. In general $\mathbf{f}$ stands for the electromagnetic and external forces for example from tides. Energy conservation implies:

$$\rho T \frac{\partial S}{\partial t} + \rho T \mathbf{v} \cdot \nabla S = \rho \epsilon - \nabla \mathbf{F} \quad (17.11)$$

where S is the entropy, ϵ is the energy generation rate per unit mass and $\mathbf{F}$ is the flux. We also require specifying the energy transport mechanism, which is the same as we have discussed in previous lectures (either radiative or convective).

We can now consider linear oscillation modes. The idea here is to consider small perturbations to 1D spherically symmetric stellar models in hydrostatic equilibrium. We therefore have an equilibrium solution to the standard time-independent equations (e.g., from running a MESA model), that can be expressed as $m_0(r)$, $P_0(r)$, $L_0(r)$, $T_0(r)$, etc. We assume that the oscillations cause 3D periodic deviations from equilibrium with small amplitudes that justify a linear approach. In practice, this means that we perturb the above equations by displacing a fluid element by $\delta \mathbf{r}$ about the equilibrium solution at position $\mathbf{r}_0$. For example a perturbation in the pressure would be calculated as:

$$\delta P(\mathbf{r}) = P(\mathbf{r}_0 + \delta \mathbf{r}) - P_0(\mathbf{r}_0) = P_0(\mathbf{r}_0) + \delta \mathbf{r} \cdot \nabla P_0 - P_0(\mathbf{r}_0) \quad (17.12)$$

The linearized equations can then be obtained by inserting expressions like Eqn. 17.12 into the conservation equations, subtracting off the equilibrium solution, and only keeping terms of linear order in the perturbation.

As we are dealing with spherical objects, the natural coordinate system is polar coordinates (r, θ, ϕ) . We can also separate the displacement vector $\delta \mathbf{r}$ into radial and horizontal components: $\delta \mathbf{r} = \xi_r \mathbf{a}_r = \xi_h$ where $\mathbf{a}_r$ is a unit vector directed radially outward. Solutions to these equations lead to eigenfrequencies which correspond to spherical modes of oscillation. We can express the displacement vector ξ as:

$$\xi(r, \theta, \phi, t) = [(\xi_r \mathbf{a}_r + \xi_h \nabla_h) Y_l^m(\theta, \phi)] \exp(-i\omega t) \quad (17.13)$$

where Y_l^m are the usual spherical harmonics, n is the radial order, l is the mode degree, m is the azimuthal order, and ω is the frequency. An equivalent, perhaps more transparent version of the above is:

$$\xi(r, \theta, \phi, t) = \left[\xi_r Y_l^m \mathbf{a}_r + \xi_h \left(\frac{\partial Y_l^m}{\partial \theta} \mathbf{a}_\theta + \frac{1}{\sin \theta} \frac{\partial Y_l^m}{\partial \phi} \mathbf{a}_\phi \right) \right] \exp(-i\omega t) \quad (17.14)$$

In the limit of spherical symmetry, the perturbed equations do not depend on m . This is no longer the case when the Coriolis force is included for rotating stars. In many analyses one

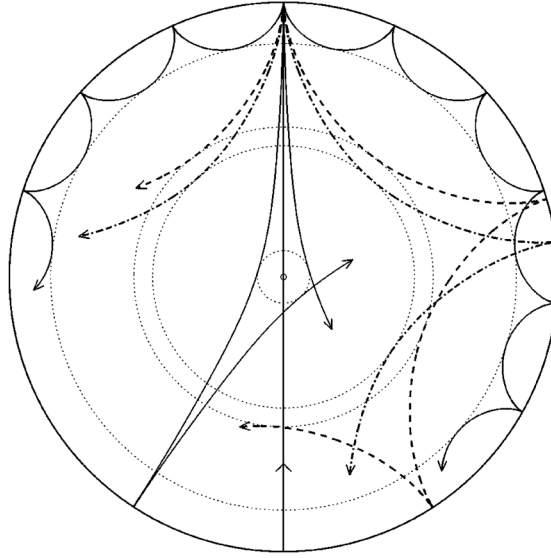


Figure 17.1: Propagation of rays of sound. The ray paths are bent by the increase in sound speed with depth (an example of refraction). Notice that rays of different angular degree reach different depths. From Christensen-Dalsgaard (2002).

makes the adiabatic approximation, which allows us to ignore perturbations in the entropy. If we make one further assumption, that the perturbations in the potential can be ignored (this is called the Cowling approximation), then the pulsation equations can be written as a single second-order differential equation for ξ_r :

$$\frac{d^2 \xi_r}{dr^2} = -\frac{\omega^2}{c_s^2} \left(\frac{N^2}{\omega^2} - 1 \right) \left(\frac{S_l^2}{\omega^2} - 1 \right) \xi_r. \quad (17.15)$$

where N is the Brunt-Väisälä frequency, ω is the frequency of the oscillation, and S_l is the characteristic acoustic frequency (also known as the Lamb frequency):

$$S_l^2 \equiv \frac{l(l+1)c_s^2}{r^2} \quad (17.16)$$

Equation 17.15 provides insight into what kinds of waves (modes) will be excited and in what regions of the star. Modes that are not oscillatory will be damped and are referred to as evanescent waves.

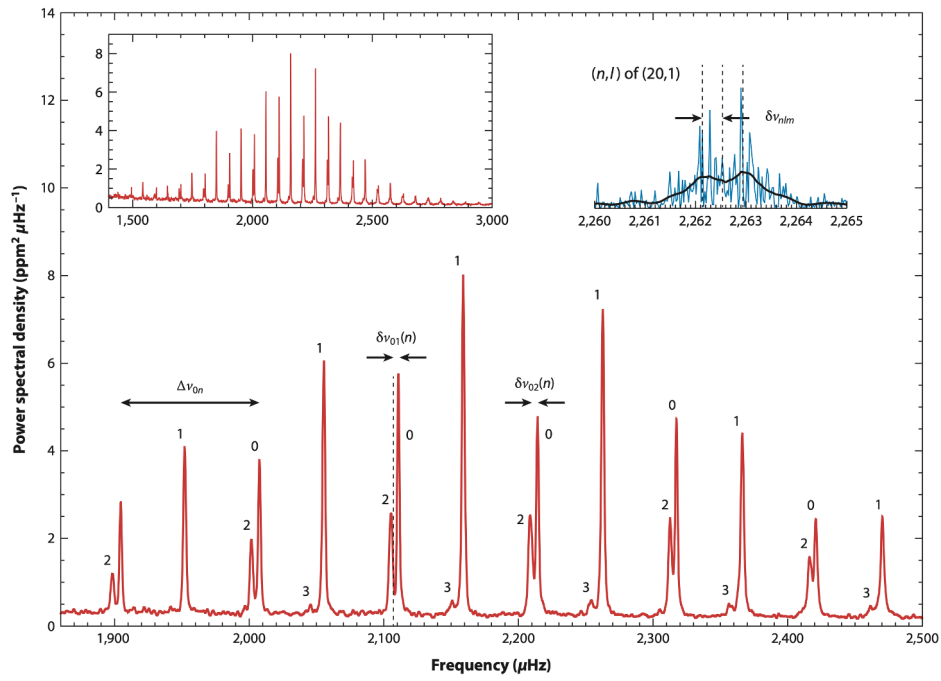


Figure 17.2: Power spectrum of the G-type main sequence star 16 Cyg A. The inset upper right panel shows a larger range of frequencies, where the maximum frequency is clearly visible at ≈ 2200 μHz . From Chaplin & Miglio (2013).

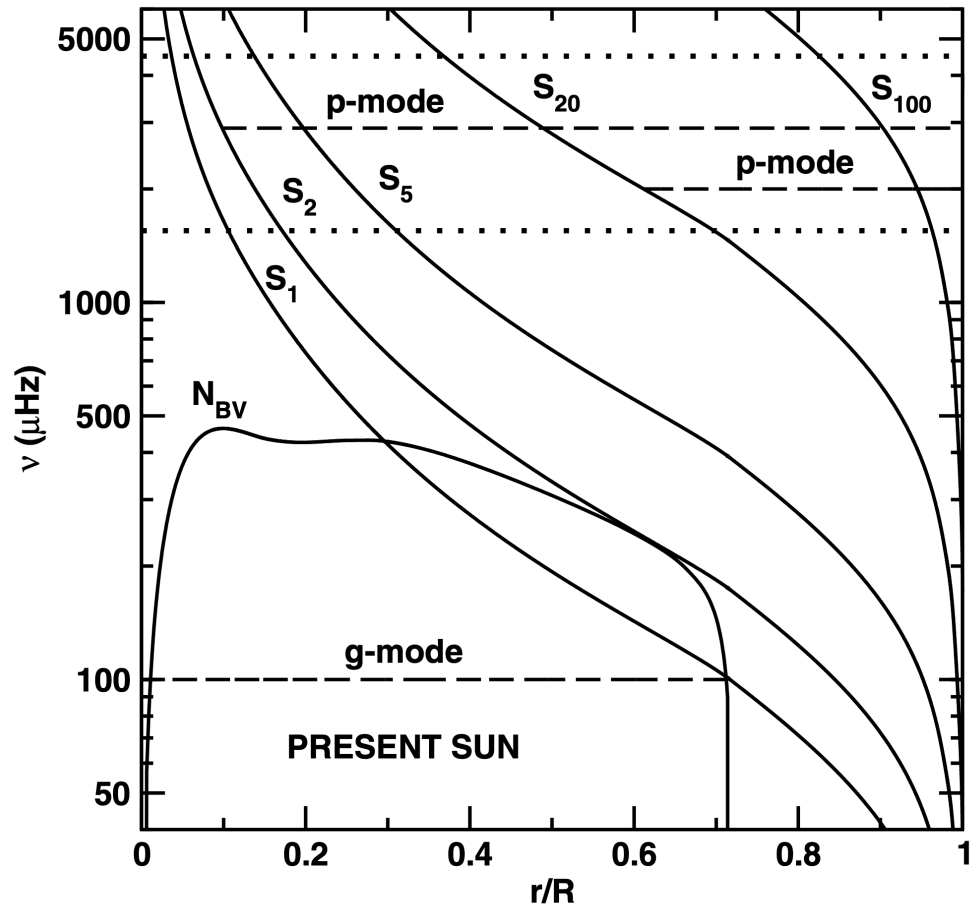


Figure 17.3: Propagation diagram for the present-day Sun. g -modes are confined to the region labeled N_{BV} , while p -modes are confined between the surface and the Lamb frequency S_L . Notice that g -modes in the Sun are confined to the core and hence have not been observed. The frequency domain corresponding to solar-like oscillations is marked by the dotted lines. From Stasinska et al. (2012).

Chapter 18

White Dwarfs

White dwarfs play an important role in astronomy. They are the end points of stellar evolution for the vast majority of stars in the Universe (all stars with $M < 8M_{\odot}$). Most of the metals in the Universe are today contained within white dwarfs. Their explosions are believed to give rise to supernovae Type Ia, the standard candles that enabled the discovery of an accelerating universe. In this lecture we will discuss some of the observational background, the mass-radius relation, internal structure, and cooling of white dwarfs.

The first white dwarf was discovered in the triple star system 40 Eridani. In 1910, Russell, Pickering, and Fleming discovered that Eri B was of spectral type A and yet was much fainter than its companion (Eri A, which is a K star). This was a major puzzle in light of what was then known about the main sequence (cooler stars should be fainter than hotter stars). Sirius B was the next object identified as a white dwarf. Importantly, the Sirius system is only 2.6 pc away, enabling detailed observations. The binary orbit of Sirius allowed measurement of the mass, and estimates of the luminosity and temperature implied a very small radius: $\approx 0.01R_{\odot}$! Such a high surface gravity led Eddington (yes, the same Eddington) to predict a gravitational redshift from Sirius B due to General Relativity effects. The gravitational redshift was measured for Sirius B in 1925. In 1926 Ralph Fowler applied ideas from quantum mechanics to deduce that electron degeneracy pressure could maintain hydrostatic equilibrium in white dwarfs. In 1931 Subrahmanyan Chandrasekhar argued that white dwarfs have an upper limit to their mass of $\approx 1.4M_{\odot}$. Fowler and Chandrasekhar won the Nobel Prize in 1983 for their work on white dwarfs. To-date, $> 300,000$ white dwarfs have been discovered, mostly through *Gaia* astrometry and CMDs.

The surface temperatures of white dwarfs range widely, from $\sim 10^5$ K to ~ 4000 K. Most

white dwarfs are of spectral type B, A, or F, and as such have strong Balmer lines. Because they have such high surface gravity, their spectral lines are extremely pressure-broadened. The white dwarf spectral classes include: DA (dominated by H lines; 80%); DB (dominated by He lines; 16%), DC (continuum spectrum with no strong absorption lines; few percent); DZ (strong metal lines present in spectrum; few percent). As many as 25% of white dwarfs show some metal lines in their spectra. This is surprising because the diffusion timescale for elements to sink out of the atmosphere is very short ($\lesssim 10^6$ yr), so one might have expected no heavy elements to remain at the surface. Remarkably, it appears that the resolution to this puzzle is the occasional accretion of rocky material from disintegrating planetesimals (the Roche radius of a rocky planet orbiting a white dwarf is $1R_\odot$).

As we have seen in previous lectures, stars with initial masses in the range $\approx 0.8 - 8M_\odot$ will eventually leave behind a carbon-oxygen (CO) white dwarf. Lower mass stars will not undergo the helium flash and will instead leave behind a helium white dwarf. The lifetimes of such low-mass stars are longer than the age of the universe, so the only way to produce helium white dwarfs is via binary star evolution with mass stripping.

Let's begin by deriving the mass-radius relation for white dwarfs supported by non-relativistic degeneracy and the upper mass limit of white dwarfs.

Non-relativistic degenerate electrons have an Equation of State (EOS) of the form:

$$P_{\text{NR,e}} \approx 10^{13} (\rho/\mu_e)^{5/3} \quad \text{dyne cm}^{-2} \quad (18.1)$$

The ultra-relativistic equivalent is

$$P_{\text{UR,e}} \approx 10^{15} (\rho/\mu_e)^{4/3} \quad \text{dyne cm}^{-2} \quad (18.2)$$

If this source of pressure can be balanced against gravity via the hydrostatic equilibrium (HSE) condition:

$$\frac{dP}{dr} = -\frac{Gm(r)\rho}{r^2} \quad (18.3)$$

then the object will be stable. If we make the usual dimensional arguments we can re-cast the HSE condition as:

$$P \sim GM^{2/3}\rho^{4/3} \quad (18.4)$$

setting this equal to the non-relativistic degenerate electron EOS leads to:

$$R \approx 2 \times 10^9 \left(\frac{0.1M_\odot}{M} \right)^{1/3} \quad \text{cm} \quad (18.5)$$

which is the familiar result that the radius decreases as the mass increases for objects supported by non-relativistic degeneracy pressure. For comparison note that the radius of earth is $R_{\oplus} = 6.4 \times 10^8$ cm — white dwarfs have radii comparable to that of the Earth while having masses comparable to that of the Sun. With a radius $\approx 10^{-2}R_{\odot}$, white dwarfs have densities of $\approx 10^6\rho_{\odot}$. Surface gravity scales as MR^{-2} , so a white dwarf will have a surface gravity 10^4 times larger than the Sun; $(\log g)_{\odot} \approx 4.5$, so $(\log g)_{\text{WD}} \approx 8.5$.

Relativistic effects become important when the Fermi momentum is comparable to $m_e c$:

$$m_e c = p_F = \left(\frac{3h^3 n_e}{8\pi} \right)^{1/3} \quad (18.6)$$

which corresponds to a critical density of $\rho_{\text{crit}} \approx 2 \times 10^6$ g cm $^{-3}$. Near this density we can expect the EOS to take the form of the relativistic degenerate case above.

Inserting the ultra-relativistic limit degeneracy pressure into the HSE equation leads to: $GM^{2/3}\rho^{4/3} = 10^{15}(\rho/\mu_e)^{4/3}$. This equation implies $M = \text{constant}$. Working out the details leads to the following expression:

$$M_{\text{Ch}} \approx 1.4 \left(\frac{2}{\mu_e} \right)^2 M_{\odot} \quad (18.7)$$

which is known as the **Chandrasekhar mass limit**. Above this mass, electron degeneracy pressure cannot support the object against its own self-gravity. Another way to understand this limit is the following: the mass-radius relation for white dwarfs implies decreasing radius with increasing mass. In the NR limit the radius decreases as $M^{-1/3}$; when relativistic effects become important the mass dependence becomes steeper (see Figure 18.1). At M_{Ch} the radius goes to zero. You will work through the implications of this behavior in HW6.

We are now going to derive the evolutionary behavior of white dwarfs.

During the AGB, low mass stars have a hot isothermal degenerate core with $T_c \sim 10^8$ K. The interior is isothermal due to efficient conductive heat transfer within the degenerate electron interior (see the supplementary material below for additional discussion of conduction in white dwarfs). This is the “initial condition” for a newly formed white dwarf.

Within a white dwarf, the ions are initially close to an ideal gas throughout the star. For the electrons, we expect them to be degenerate in the core and (close to) ideal at the surface. We can compute at which location this transition occurs by noting that both forms of pressure will be of comparable magnitude where $E_F \approx kT$ where E_F is the Fermi Energy:

$$E_F = \frac{1}{2m_e} \left(\frac{3h^3 n_e}{8\pi} \right)^{2/3} \quad (18.8)$$

which implies that the transition density will be $\rho_t \approx 1.2 \times 10^4 T_8^{3/2} \text{ g cm}^{-3}$ where T_8 is the temperature in units of 10^8 K . This is one equation with two unknowns. To make further progress we need a second equation, which we can derive by appealing to hydrostatic equilibrium.

In the outer layers of the white dwarf where $M_{\text{outer}} \ll M$ we can approximate the structure as a geometrically thin plane parallel atmosphere. This means we can assume that the surface gravity g is constant. Typical values for a white dwarf are $\log g \approx 8$ (in cgs units). Under these assumptions the pressure is:

$$P = \frac{g M_{\text{outer}}}{4\pi R^2} = \frac{\rho k T}{\mu m_p} \quad (18.9)$$

Combining this with the white dwarf mass-radius relation and the earlier expression for ρ_t leads to the following:

$$M_{\text{outer}} = 10^{-3} T_8^{5/2} M_{\odot} \quad (18.10)$$

where T_8 is the temperature of the isothermal degenerate core in units of 10^8 K . The implication is that $\approx 99.9\%$ of the white dwarf is degenerate (by mass), and therefore the vast majority of the object is (nearly) isothermal. But we know that the surface of the white dwarf is much cooler than $\sim 10^8 \text{ K}$; the extremely large temperature drop between core and surface must therefore occur in the very thin envelope. The transfer of radiation through the white dwarf envelope therefore plays a key role in the cooling and evolution of the white dwarf. Below we will assume that radiation is the primary means of flux transport through the envelope in order to estimate the cooling rate of white dwarfs. This is a good assumption except for the very coolest white dwarfs, which develop convective envelopes. In the preceding derivations we have defined the core-envelope boundary as the boundary between degenerate and ideal gas conditions. In the following discussion we assume that this boundary also marks the transition from conductive energy transport (in the interior) to radiative energy transport (in the envelope).

If we now assume the flux, F , passing through this outer layer is constant and the flux transport is via radiation then we have:

$$F = \frac{1}{3} \frac{c}{\kappa \rho} \frac{d(aT^4)}{dr} \quad (18.11)$$

We can use the HSE equation to replace ρdr with dP/g :

$$\frac{3F\kappa}{c} \frac{dP}{g} = 4aT^3 dT \quad (18.12)$$

Assuming Kramers' opacity ($\kappa = \kappa_0 \rho T^{-7/2}$) and an ideal gas EOS to replace ρ with P leads to:

$$\left[\frac{3F}{4c} \frac{\mu m_H \kappa_0}{kga} \right] P dP = T^{7.5} dT \quad (18.13)$$

If we now integrate this equation from the base of the envelope (the transition point between degenerate and ideal EOS that we derived above) to the surface and set $F = L/4\pi R^2$, then we arrive at **Mestel's Cooling Law**:

$$L = 2.5 L_\odot \left(\frac{M}{0.6 M_\odot} \right) \left(\frac{T_c}{10^8 \text{ K}} \right)^{7/2} \quad (18.14)$$

Nearly all of the energy in the white dwarf is stored in the thermal content of the ions, which obeys an ideal EOS:

$$E_{\text{th}} = \frac{3}{2} N_i k T_c = 1.3 \times 10^{48} T_8 \left(\frac{M}{0.6 M_\odot} \right) \text{ erg} \quad (18.15)$$

Since there is no nuclear fusion (no *production* of energy), the luminosity can be expressed as:

$$L = - \frac{dE_{\text{th}}}{dt} \quad (18.16)$$

If we equate this expression with Eqn. 18.14 then we find that the mass-dependence cancels. We are left with

$$\frac{T_8^{7/2}}{4 \times 10^6 \text{ yr}} = - \frac{dT_8}{dt} \quad (18.17)$$

after a few e-foldings the core temperature evolves as:

$$T_8 = \left(\frac{1.6 \times 10^6 \text{ yr}}{t} \right)^{2/5} \quad (18.18)$$

which, via Eqn. 18.14 implies $L \propto t^{-7/5}$.

We also have, via Stefan-Boltzmann: $L = 4\pi R^2 \sigma T_{\text{eff}}^4$. Combining this with the mass-radius relation ($R \propto M^{-1/3}$) gives relations between M , L , and T_{eff} .

We can measure the mass of white dwarfs from their spectra with a reasonably high degree of accuracy. The Balmer lines are significantly pressure-broadened, enabling a direct measurement of the surface pressure. We know from the plane-parallel HSE equation that $P \propto g/\kappa$, meaning that an estimate of the pressure provides a measurement of M/R^2 . We also know that white dwarfs obey a well-understood mass-radius relation, which allows us to obtain a simple relation between surface pressure and mass.

So far we have ignored the fact that the ions interact with each other via the Coulomb force. We can estimate when this will be important by defining the **Coulomb parameter**:

$$\Gamma \equiv \frac{Z^2 e^2}{r_0 k T} \quad (18.19)$$

where r_0 is the average ion spacing and Z is the charge. Notice that when $\Gamma = 1$ we have $Z^2 e^2 / r_0 = kT$, i.e., the electrostatic potential energy is equal to the thermal energy per degree of freedom. When $\Gamma > 1$ we can expect that the ions will interact strongly via Coulomb interactions. In a liquid (when $\Gamma \gtrsim 1$) $E_{\text{th}} = 3kT$ whereas in a gas $E_{\text{th}} = \frac{3}{2}kT$, owing to the change in the number of degrees of freedom. At $\Gamma \approx 180$ the ions within the white dwarf **crystallize**. The crystallization process releases a *latent heat* due to the first-order phase transition. This release of heat delays the cooling process compared to the simple derivation above (by about a factor of two). This phase transition occurs several Gyr after the formation of the white dwarf, and after it occurs basically all of the white dwarf is crystallized. The degenerate electrons still dominate the pressure, even when the ions are crystallized. An additional complication is the sedimentation of “impurities” in the white dwarf, including ^{22}Ne and ^{56}Fe , which releases gravitational energy due to chemical differentiation and hence is an additional source of energy affecting the cooling process. As the white dwarf cools further, eventually the crystallized white dwarf enters the Debye regime in which the heat capacity drops rapidly according to $c_V \propto T^3$. This results in accelerated cooling beyond the standard Mestel model and is known as Debye cooling.

Supplementary Material:

Here we show that heat transfer by conduction is more efficient than radiation in white dwarf interiors.

Recall from Lecture 7 that heat transfer by a generic diffuse process can be described by Fick’s law:

$$F = -D \frac{dT}{dr} \quad (18.20)$$

where F is the flux and D is a diffusion coefficient. In the case of conduction via electrons we have:

$$F_{\text{cond}} = -\frac{4ac}{3\kappa_{\text{cond}}\rho} T^3 \frac{dT}{dr} \quad (18.21)$$

where

$$\kappa_{\text{cond}} \equiv \frac{4acT^3}{3D\rho} \quad (18.22)$$

is the “conductive opacity”. We have defined κ_{cond} in this way so that we can write the total opacity as the harmonic sum of the radiative and conductive opacities:

$$\frac{1}{\kappa_{\text{tot}}} = \frac{1}{\kappa_{\text{rad}}} + \frac{1}{\kappa_{\text{cond}}} \quad (18.23)$$

and then the total flux is simply:

$$F_{\text{tot}} = -\frac{4ac}{3\kappa_{\text{tot}}\rho} T^3 \frac{dT}{dr} \quad (18.24)$$

We can see from the above equations that whatever opacity is *smaller* is the most important for determining the heat flow. In most stellar conditions the conductive opacity is very large, which means that radiative heat flux is the dominant source of diffuse heat transfer. However, in degenerate regions the conductive opacity is small and hence conduction dominates the heat flow, as we will now show.

The diffusion coefficient can be estimated as $D \sim c_V v_e \lambda$ where c_V is the specific heat of the degenerate electrons, v_e is the typical electron velocity, and λ is the electron collisional mean free path. As we saw in HW3 (see also HKT Eqn 3.114), the specific heat of degenerate electrons can be expressed as:

$$c_V \approx \frac{8\pi^3 m_e^2 c}{3h^3} k^2 T x_F \text{ erg cm}^{-3} \text{ K}^{-1} \quad (18.25)$$

where we have assumed that x_F is small so we were able to ignore the $(1 + x_F^2)$ term in HKT Eqn 3.114. Recall that $x \propto (\rho/\mu_e)^{1/3}$.

Collisions involving degenerate electrons involves the additional constraint that the electron cannot be scattered into an already filled energy state. This implies that only electrons near the top of the Fermi sea can participate in the conduction process. We can therefore estimate the velocity entering in the diffusion coefficient via $m_e v_e \approx p_F \propto x \propto (\rho/\mu_e)^{1/3}$. The most efficient means of scattering the electrons is via Coulomb interactions with the ions. Thus: $\lambda = 1/(\sigma_C n_I)$ where n_I is the number density of ions and σ_C is the Coulomb scattering cross section. An order of magnitude estimate can be obtained by setting $\sigma_C \sim \pi s^2$ where s is the impact parameter determined by equating the electron kinetic energy and the electron-ion electrostatic potential: $m_e v_e^2 \sim Z e^2/s$. This latter condition is where we may expect significant scattering to occur. Putting these pieces together results in:

$$D \propto \frac{\mu_I}{\mu_e^2} \rho T \quad (18.26)$$

where μ_I is the ion mean molecular weight.

Putting everything together we arrive at a rough approximation to the conductive opacity:

$$\kappa_{\text{cond}} \approx 4 \times 10^{-8} \frac{\mu_e^2}{\mu_I} Z_c^2 \left(\frac{T}{\rho} \right)^2 \text{ cm}^2 \text{ g}^{-1} \quad (18.27)$$

where Z_c is the ion charge. A typical white dwarf interior has $T \approx 10^7$ K and $\rho \approx 10^6$ g cm⁻³ and is composed mostly of carbon and oxygen. The high temperature implies that the radiative opacity is electron scattering and hence $\kappa_{\text{rad}} = 0.2 \text{ cm}^2 \text{ g}^{-1}$. Our equation for the conductive opacity under these conditions yields $\kappa_{\text{cond}} \approx 5 \times 10^{-5} \text{ cm}^2 \text{ g}^{-1}$ and thus $\kappa_{\text{cond}} \ll \kappa_{\text{rad}}$. In other words, conduction is the dominant source of heat transfer in white dwarf interiors (similar arguments show that conduction is also the dominant source of heat transfer in degenerate helium cores, for example of stars on the red giant branch).

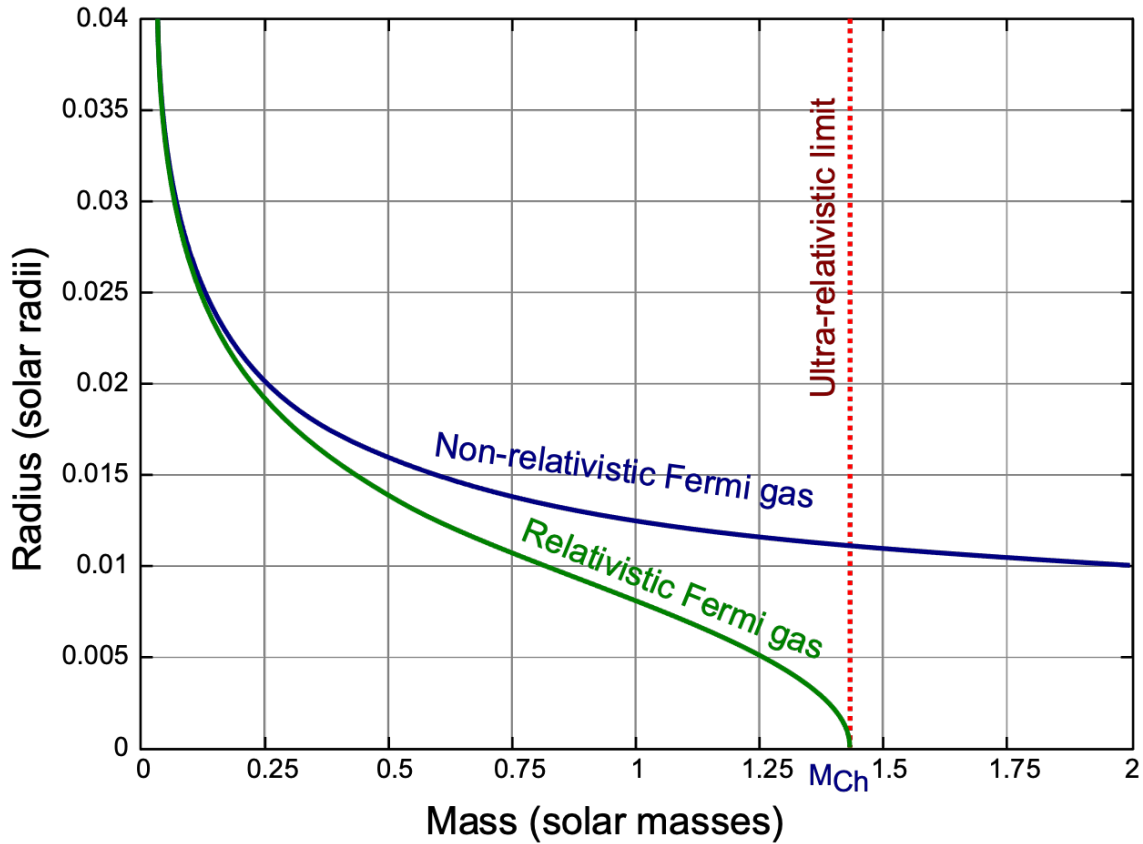


Figure 18.1: Mass-radius relation for white dwarfs. The blue line represents the non-relativistic limit ($R \sim M^{-1/3}$), while the green line represents the full Chandrasekhar theory (smoothly transitioning from non-relativistic at low densities [low masses] to ultra-relativistic at high densities [high masses]). Notice that the radius goes to zero at M_{Ch} . From Wikipedia.

Chapter 19

Neutron Stars & Black Holes

In this lecture we will discuss the physics and observable signatures of neutron stars and black holes. Together with white dwarfs, these objects are collectively known as “compact objects”.

Neutron stars arise from stars with initial masses in the range of $\approx 8 - 25M_{\odot}$ although the exact range is uncertain and depends on many parameters such as metallicity, stellar rotation, binary evolution, mass-loss, etc. Some models place the upper mass boundary as high as $40M_{\odot}$.

At sufficiently high densities ($\rho \approx 10^7 \text{ g cm}^{-3}$) nuclei will undergo inverse β -decay: $(A, Z) + e^- \rightarrow (A, Z - 1) + \nu_e$ where ν_e is an electron neutrino (recall that Z is the proton number and A is the total number of nucleons - protons plus neutrons). This will occur when the electron energy is $\sim 1.3 \text{ MeV}$, which is the rest-mass energy difference between a neutron and a proton. Eventually the nuclei are so neutron-rich that the reaction: $(A, Z) \rightarrow (A - 1, Z) + n$ occurs. This begins at a density of $\rho \approx 4 \times 10^{11} \text{ g cm}^{-3}$ and is known as “neutron drip”. Nuclei will completely dissolve into neutrons at a density of $\rho \approx 3 \times 10^{14} \text{ g cm}^{-3}$, which is the density of nuclear matter. This can be estimated by considering that the “radius” of an ordinary nucleon is $r_n \sim A^{1/3} \text{ fm}$ where $1 \text{ fm} = 1 \text{ Fermi} = 10^{-13} \text{ cm}$.

Once the core has been converted to a pure neutron gas, it will be at densities where degeneracy pressure is the dominant form of pressure. As with electrons, there will be non-relativistic and ultra-relativistic limits:

$$P_{\text{n,NR}} = K_{\text{NR}} n_n^{5/3} \quad \rho \lesssim 6 \times 10^{15} \text{ g cm}^{-3} \quad (19.1)$$

$$P_{\text{n,UR}} = K_{\text{UR}} n_n^{4/3} \quad \rho \gtrsim 6 \times 10^{15} \text{ g cm}^{-3} \quad (19.2)$$

where n_n is the number density of neutrons. In the non-relativistic limit neutron stars obey a mass-radius relation similar to white dwarfs but with a very different zero-point:

$$R_{\text{ns}} \approx 4.9(M_{\text{ns}}/M_{\odot})^{-1/3} \text{ km} \quad (19.3)$$

The size of neutron stars is comparable to that of a moderately-sized city. Observations indicate neutron star sizes closer to 10 km. Measurements of the mass-radius relation place strong constraints on the neutron star EOS.

As in the case of white dwarfs, as the density increases the EOS becomes relativistic, which would suggest an upper mass for neutron stars. However, there isn't a direct analog to the Chandrasekhar mass for white dwarfs. In the case of neutron stars, nucleons are providing the pressure *and* the gravity. Remember that in special relativity there is a concept of relativistic mass that is different (larger) than the rest mass. For relativistic neutrons the relevant mass in the hydrostatic equilibrium equation is therefore not simply $m_n n_n$. For this reason relativistic neutrons are not $n = 3$ polytropes and hence do not have an associated Chandrasekhar-like mass.

For neutron stars the density and gravity are so large that General Relativity effects matter. The full GR equation for hydrostatic equilibrium is:

$$\frac{dP}{dr} = -\left(\rho + \frac{P}{c^2}\right) \left(\frac{Gm(r)}{r^2} + \frac{4\pi GPr}{c^2}\right) \left(1 - \frac{2Gm(r)}{c^2 r}\right)^{-1} \quad (19.4)$$

this is known as the Tolman-Oppenheimer-Volkoff (TOV) equation and was derived in 1934/1939. This equation is the GR equivalent of our usual hydrostatic equilibrium equation. When combined with an EOS, this equation give the complete mechanical structure of a neutron star. Note that in the above equation the density ρ is *not* the rest-mass value ($m_n n_n$) but rather $\rho = U/c^2$ where U is the internal energy per unit volume, including rest-mass. Integrating the TOV equation with the standard degenerate EOS (Eqn 19.2) gives a maximum mass of $0.7M_{\odot}$. This limit is considerably smaller than observations, which indicate a maximum mass of $2 - 3M_{\odot}$. The use of a more accurate EOS for neutron stars produces higher maximum masses that are in better agreement with observations. Objects exceeding the maximum mass are expected to collapse to form black holes.

Neutron stars are born very hot (due to the very hot core at the end stages of massive stars), with $T \sim 10^9$ K. Like white dwarfs, they radiate away their thermal energy via photons and neutrinos. At high temperature neutrinos are the dominant coolant. The surface temperatures of neutron stars are $\sim 10^6$ K and they remain that hot for at least 10^4 yr. We can expect neutron stars to be rotating very rapidly and have high magnetic

field strengths due to the enormous collapse factors of the core – about a factor of 10^4 – and assuming approximate conservation of angular momentum and magnetic flux ($J \propto mvr$ and $B \propto r^{-2}$). Taking round numbers, if we start with a core of size 10^{11} cm rotating at 10 km s^{-1} with a magnetic field of 1 G and compress it by a factor of 10^4 we end up with an object with a magnetic field strength of 10^{10} G rotating near the speed of light. This is basically what we observe in neutron stars: very rapid rotation and very strong magnetic fields.

As long as the magnetic axis of the NS is misaligned with the rotational axis (by even a tiny amount) we will have an accelerating magnetic moment, which will radiate. The origin of the misalignment is unknown, but could be due to small asymmetries in the collapse. The dipole energy loss is given approximately by $\dot{E}_{\text{dip}} \approx R^6 B^2 / c^3$. The rotational energy of a uniformly rotating sphere is $E_{\text{rot}} = I\Omega^2/2$ where I is the moment of inertia and $\Omega = 2\pi/P$ is the spin rate and P is the rotational period. The time derivative is $\dot{E}_{\text{rot}} = -I\Omega\dot{\Omega} \propto \dot{P}/P^3$. Assuming that the radiated energy is coming at the expense of the rotational energy we can set $\dot{E}_{\text{dip}} = \dot{E}_{\text{rot}}$, which implies $B \propto \sqrt{P\dot{P}}$. We can also define a characteristic age via $\tau \equiv P/(2\dot{P})$. We can therefore estimate three important quantities: magnetic field strength, age, and luminosity, from measurement of the period and period derivative. Figure 19.1 shows a $P - \dot{P}$ diagram with lines of constant B and age. But how might we measure the spin period of neutron stars?

It turns out that the strong magnetic fields give rise to emitting beams of electromagnetic radiation along the direction of the magnetic poles. If the beam intersects with our line of sight, then we observe the radiation as a **pulsar**. The opening angle is relatively small, so one can imagine that the observed number of pulsars is smaller than the intrinsic number within the Galaxy (models suggest a factor of ~ 5 due to the beam geometry). Pulsars are observed in the radio, although this radio emission is a tiny fraction of the total emitted energy. Most of the energy loss is in the form of high energy radiation (gamma rays) and a particle wind (known as a pulsar wind). The first pulsar was discovered by Jocelyn Bell and Antony Hewish in 1967 via repeating radio signals.

Pulsars are extremely accurate clocks (more accurate than terrestrial atomic clocks). A neutron star binary (in which one member is a pulsar) was used to provide the first indirect detection of gravitational waves via the decay of the binary orbit (this binary, PSR 1913+16, is known as the Hulse-Taylor binary). Large networks of pulsars are also now being used to detect gravitational waves (known as pulsar timing arrays; more on this in the next lecture). Pulsars spin down over time, and it is expected that after 10 – 100 million years the pulsar will “turn off”. For this reason only a tiny fraction of neutron stars in the Galaxy are believed

to be pulsars.

Neutron stars with $B > 10^{13}$ G are known as **magnetars**. There are 23 known magnetars, most of them are in our Galaxy. Their very strong magnetic fields decay on a short timescale of $\sim 10,000$ yr, explaining their relative rarity. The strong fields give rise to lots of exotic high energy phenomena, including soft gamma ray repeaters (SGRs) and anomalous X-ray pulsars (AXPs). They may also give rise to fast radio bursts (FRBs). The origin and properties of magnetars are active areas of research.

In Figure 19.1 there is a distinct set of pulsars with very short periods and very small spin-down rates (the lower left region of the plot). These are the millisecond pulsars (MSPs), nearly all of which are found in binary systems. The formation of MSPs is believed to be related to accretion onto a neutron star from a companion. This accretion spins the neutron star up to very high spin rates. For this reason MSPs are sometimes called recycled pulsars. For unknown reasons MSPs have very weak magnetic fields. The weak fields imply very low spin-down rates.

Black holes are objects so compact that nothing can escape from them, not even light. Such objects have been imagined since at least the eighteenth century, as it is straightforward to imagine an object whose escape velocity, $v_{\text{esc}}^2 = 2GM/R$, is greater than c^2 . Of course, at such large velocities and gravitational fields, Newtonian gravity breaks down and so we had to await the arrival of General Relativity before real progress could be made.

It is important to remember that, owing to the fact that black holes are by definition dark, it has been difficult to definitively prove that they exist. The most direct evidence that such objects must exist come from mass constraints via orbital dynamics, whether in the center of our Galaxy (in the case of supermassive black holes) or binary star systems in the case of stellar-mass black holes. In the latter case, the basic idea is to model the radial velocity curve, estimate the mass of the luminous star, and find invisible companions with masses $\gtrsim 3M_{\odot}$ (safely above the upper limit for neutron stars). Additional indirect evidence comes from the light emitted by accretion disks surrounding black holes, and more recently, via the image of the black hole event horizon by the EHT collaboration based at the CfA and gravitational wave events from LIGO.

Stellar-mass black holes have masses of $\sim 3 - 100M_{\odot}$ and are believed to have formed from the collapse of massive stars (with initial masses $\gtrsim 40M_{\odot}$). Supermassive black holes (SMBH) have masses of $\sim 10^4 - 10^9 M_{\odot}$ and grow primarily through the accretion of gas from their surroundings. The initial seeds of SMBHs are not known, but could be stellar-mass

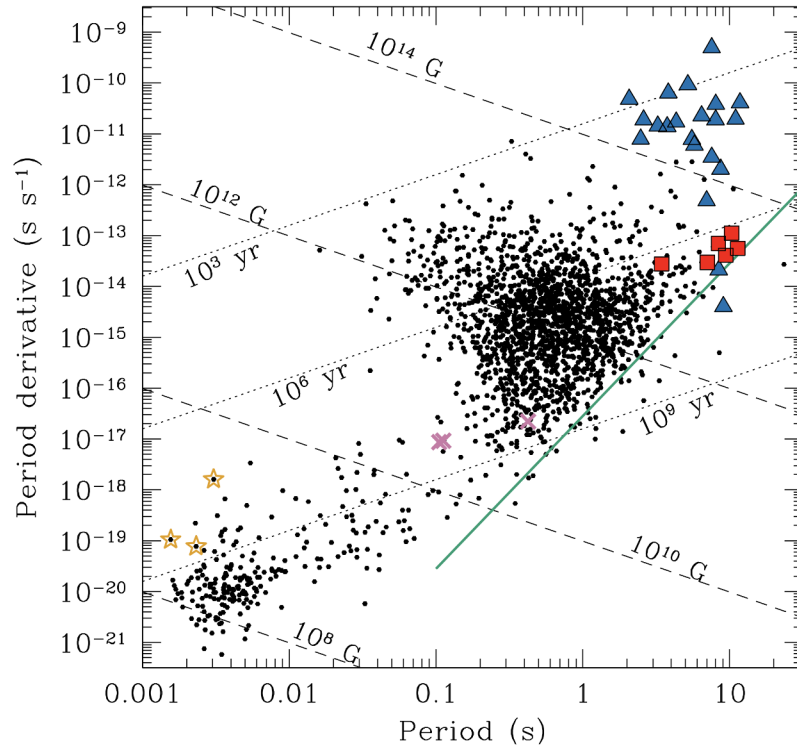


Figure 19.1: $P - \dot{P}$ diagram for pulsars. Lines of constant magnetic field strength and age are shown. Triangles indicate known magnetars; points in the lower left are millisecond pulsars (“recycled” pulsars, nearly all of which are in binary systems). The green line is the theoretical pulsar death line, below which pulsars are expected to cease producing radio emission. From Ray (2019).

black holes or direct-collapse of large gas clouds in the early universe (the latter are known as direct-collapse black holes). There is a third category of intermediate-mass black holes (IMBHs). There have been many reported detections of IMBHs, but the reality of such objects is still under active debate.

The **Schwarzschild radius**, R_s , is the event horizon of a non-rotating black hole, within which nothing will escape. It is the location at which the escape speed is equal to the speed of light, hence $R_s = \frac{2GM}{c^2}$. BH spin is characterized by a dimensionless spin parameter:

$$a \equiv \frac{cJ}{GM^2} \quad (19.5)$$

where J is the BH angular momentum. Mass and spin are the two primary characteristics of BHs.

The location at which stable orbits can be maintained is known as the innermost stable circular orbit, or ISCO, and it depends on the rotation rate of the black hole. This fact allows

the indirect estimate of the spin of black holes. The ISCO is $R_{\text{ISCO}} = 3R_s$ for non-spinning BHs and for a maximally spinning BH the ISCO of prograde material is $R_{\text{ISCO}} = 0.5R_s$. Assuming that an accretion disk extends to the ISCO, a measurement of the maximum velocity of orbiting material provides a measurement of BH spin. Another approach is to adopt a model for the accretion disk structure, including the temperature profile. The maximum temperature is then a probe of the inner disk and hence the ISCO.

BHs grow either through mergers with other compact objects or via an accretion disk. In the latter case, we can estimate the maximum growth rate by setting the Eddington luminosity equal to accretion luminosity. The latter can be parameterized as

$$L_{\text{acc}} = \epsilon \dot{M} c^2 \quad (19.6)$$

where ϵ is the radiative efficiency and quantifies the fraction of the rest mass energy that is transferred into radiation. Setting $L_{\text{Edd}} = L_{\text{acc}}$:

$$4\pi c G M m_H / \sigma_T = \epsilon \dot{M} c^2 \quad (19.7)$$

which implies:

$$\int d \ln M = \frac{4\pi G m_H}{\epsilon \sigma_T c} \int dt \quad (19.8)$$

or

$$M_{\text{BH}}(t) = M_0 e^{t/t_S} \quad (19.9)$$

where t_S is the Salpeter time:

$$t_S \equiv \frac{\epsilon \sigma_T c}{4\pi G m_H} = 4.5 \epsilon_{0.1} \times 10^7 \text{ yr} \quad (19.10)$$

where $\epsilon_{0.1} = \epsilon/0.1$. Black holes growing at the Eddington limit grow exponentially with a timescale set by the Salpeter time.

The existence of luminous quasars at $z > 6$ presents a puzzle in this context. Quasars are believed to be accreting supermassive BHs, and models suggest that the detections at $z > 6$ have BH masses of $\sim 10^9 M_\odot$. If we start with a stellar mass BH seed of mass $10 M_\odot$, it would have to accrete at L_{Edd} for ≈ 800 Myr in order to grow to $\sim 10^9 M_\odot$. But at $z = 6$ the universe is only 900 Myr old – so the puzzle is how to grow supermassive BHs rapidly in the early universe. One solution is super-Eddington accretion, and another is to appeal to more massive seed BHs, e.g., via direct collapse of massive gas clouds, which would create seeds as massive as $\sim 10^3 M_\odot$. This topic is a very active area of research.

Chapter 20

Supernovae I

Observationally, a supernova is a powerful and luminous transient event. Supernovae (SNe) have been observed since antiquity, including eight that have been documented in the pre-telescopic historical record. In terms of the underlying physical mechanism, there are two basic types: **core-collapse** of a massive star and **thermonuclear detonation** of a white dwarf. The latter are classified observationally as Type Ia, while the former have a wide range of types, including Type II, Type Ib, Type Ic, etc.

As we discussed in earlier lectures, massive stars are able to fuse elements through silicon burning, after which point fusion is unable to provide any meaningful amount of thermal pressure to balance gravity. The result is the collapse of the core. The outcome of this collapse is very complicated and likely gives rise to a diversity of outcomes depending on the initial stellar mass, rotation rate, mass-loss, binarity, and metallicity. Outcomes include explosion+NS, explosion+BH, explosion with no remnant, and implosion+BH.

In the mass range $\sim 8 - 10M_{\odot}$ the core will burn beyond helium but not all the way to silicon. An OMgNe degenerate core develops, and if the density is sufficiently high ($\sim 10^{10} \text{ g cm}^{-3}$) then oxygen will ignite. Because the core is degenerate, the ignition will be *explosive* and lead to an **electron-capture SNe**. These objects are so-named because the collapse of the core is triggered by electron capture on ^{24}Mg and ^{30}Ne , which remove pressure support that was being provided by the electrons. It is unclear what if any remnant is left behind, but possibilities include a neutron star and a bound ONeFe white dwarf.

Above $\sim 10M_{\odot}$ the stars all undergo fusion reactions all the way to an iron core of mass $\approx 1 - 2M_{\odot}$. As we saw in previous lectures, beyond carbon burning the Coulomb barriers are very high, requiring very high temperatures to overcome. At the same time, neutrino losses scale as a relatively high power of temperature, which means that much of the energy

generated by fusion is lost (carried away) by the neutrinos. The timescale for the last few stages of nuclear burning are therefore very short (days to weeks). An important consequence of this very short nuclear timescale in the core is that the outer layers (in which the thermal timescales are much longer) are essentially frozen in time as the core proceeds beyond carbon burning. For example, in a $15M_{\odot}$ star silicon burning lasts 18 days and the collapse of the iron core occurs in ~ 1 second, which is much shorter than even the dynamical timescale of the envelope (hours to days).

The iron core collapses until $T > 10^{10}$ K. At such high temperatures photodissociation breaks apart heavy nuclei, resulting in reactions such as $^{56}\text{Fe} + \gamma \rightarrow 13^4\text{He} + 4\text{n}$. This reaction is *endothermic*, so does not halt the collapse (we are glossing over a lot of complicated nuclear physics here). The collapse continues until nuclear densities are reached, e.g., $\rho \sim 2 \times 10^{14}$ g cm $^{-3}$. This final collapse occurs on millisecond timescales.

The collapse creates a proto-neutron star for cores below the Oppenheimer-Volkoff limit (the limit above which neutron degeneracy pressure cannot balance gravity), which is $\approx 2M_{\odot}$. The outer layers (above the proto-NS) are in-falling supersonically and eventually collide with the surface of the NS. This collision, often referred to in the literature as a “bounce”, generates a shock wave that travels outward. One can think of this bounce as piston acting on the inner layers of the collapsing star (1D models rely on this piston model to simulate the subsequent behavior). It was once thought that this outwardly-propagating shock wave could power the explosion of the star. However, it is now known, based on detailed simulations, that the shock *stalls* because photodissociation and electron capture drains energy from the shock. In addition, ram pressure (ρv^2) from the falling outer layers of the star is considerable, and presents an additional obstacle for the shock to overcome. A major industry has developed in the field over the past 40 years devoted to figuring out how to “revive” the stalled shock.

There is another problem: while the shock is stalled, accretion onto the proto-NS is occurring at a rate of $\sim 0.1M_{\odot} \text{ s}^{-1}$! If the accretion continues even for a second, a BH will form (as the neutron star will exceed the Oppenheimer-Volkoff limit). This means that the material above the proto-NS must be ejected very quickly if a NS is to survive.

The proto-NS will radiate $\sim 10^{53}$ erg of neutrinos in a few seconds. These neutrinos come from the neutronization of the core (electron captures). The energy source ultimately comes from the binding energy of the contracting proto-NS. With a mass of $\approx 1.4M_{\odot}$ and $R = 20$ km we have $E = 3 \times 10^{53}$ erg. The typical kinetic energy of a core-collapse SNe is $\sim 10^{51}$ erg (assuming $V \sim 10^4$ km s $^{-1}$ and $M_{\text{env}} \sim 10M_{\odot}$), and the typical energy in photons is only 10^{49} erg. There therefore need to tap into only $\sim 1\%$ of the neutrino energy in order to

explode the star. It turns out that, at the densities of interest, the shock is optically-thick to neutrinos. Current models favor this “neutrino-driven” shock-wave as the means to revive the stalled shock. The process is inherently multi-dimensional, which necessitates 3D models and makes simulations extremely computationally expensive. There are lots of complicated instabilities, uncertainties in the EOS, particle physics data, etc.

During the collapse, the envelope is also falling inward and experiences an increase in temperature and density, which results in dramatically higher fusion rates. This rapid injection of energy into the envelope aides in the explosion.

If the accretion onto the proto-NS proceeds for a sufficiently long period of time, a BH will form. The inner core is essentially decoupled from the outer layers, so at this point, whether a BH forms or not is unlikely to affect the subsequent SN event.

At the highest initial masses ($\gtrsim 130M_{\odot}$) the helium core is convective and radiation pressure dominates the total pressure, meaning that the adiabatic index is close to $\gamma = 4/3$. After helium burning the carbon-oxygen core contracts and ignites carbon burning. At about $T \approx 10^9$ K electron-positron pairs are produced in large numbers ($\gamma + \gamma \leftrightarrow e^- + e^+$). Pair production removes the dominant pressure source (radiation) from the core and results in contraction. Contraction results in an increase in temperature, accelerating the pair creation rate. This is a runaway process. The core collapses, and in the process oxygen ignites explosively. When the energy liberated from oxygen fusion is large enough the collapse ceases and is reversed into an explosion. This thermonuclear explosion is known as a **pair instability supernova** (PISN), which is expected to completely destroy the star.

At even higher initial masses (e.g., $\gtrsim 260M_{\odot}$), the energy generated from explosive oxygen burning is insufficient to overcome the binding energy of the star. The fate of such objects is unclear, but current models predict that they collapse to black holes without producing a luminous supernova.

Gamma-ray bursts (GRBs) are extremely energetic transient events in gamma rays. They were first discovered in 1967 by the Vela satellites, which were designed by the U.S. to detect nuclear weapons tests. Very little was known about the sources, although the sources were apparently isotropic across the sky, suggesting an extragalactic origin. A major breakthrough came in 1997 when afterglows from GRBs were discovered in X-rays, optical, and radio waves. This enabled precise localization of the burst and an association with a galaxy and a measurement of the redshift and hence distance. This conclusively demonstrated a cosmological origin of GRBs and therefore (because of their great distance) the enormity of

the associated energy ($\sim 10^{51}$ erg).

The gamma ray emission is believed to originate from a collimated narrow jet with opening angles in the range of $2 - 20^\circ$ and Lorentz factors of $\Gamma \equiv (1 - v^2/c^2)^{-1/2} \sim 100$. Because we can only observe a GRB when the jet is pointed at us, we expect the true GRB rate to be much higher than the observed rate (as in the case of pulsars).

Data from the Compton Gamma-Ray Observatory (1991-2000) revealed that GRBs can be divided into two broad categories: **short-hard** and **long-soft**, where short/long refers to the timescale of the event and shows a separation at ≈ 2 s (hard/soft refers to the shape of their non-thermal energy spectrum, with “hard” referring to relatively more high-energy emission compared to “soft”).

Long duration GRBs, which are transient on timescales of ~ 20 s, have long been associated with star forming galaxies, and have been definitively linked to optical SNe in a few cases. Even after accounting for the opening angle effect, not all long duration GRBs are associated with SNe (at least one has been associated with a kilonova!). The favored model for this class of GRB is the **collapsar** model. In this model, a rotating massive star collapses into a black hole and produces an accretion disk. A rotating black hole fed by an accretion disk tends to produce a jet of the right properties to match the observational constraints. This model requires high rotation rates at the end of a star’s life, as well as large helium cores. Both of these are easier to achieve with mass-loss rates are lower, i.e., at low metallicity. So we might expect this class of GRBs to be more common at low metallicity (e.g., low-mass galaxies and/or higher redshift).

For the short duration GRBs, the typical timescale of ≈ 0.2 s implies a very small physical scale of order an Earth radius. It took longer to sort out the origin of these events because counterparts at other wavelengths were difficult to find. The first few X-ray associations placed short duration GRBs in elliptical galaxies with no active star formation. This ruled out an association with massive stars. In fact, these events were long believed to be associated with binary neutron star mergers based on very general energetic and physical size arguments. The definitive moment came in 2017 when a GRB was associated with the gravitational wave event GW170817 and a subsequent optical transient event with an unusual light curve. All of these observations fit beautifully with models of binary neutron star mergers. The optical event is known as a kilonova, as the peak brightness was predicted to be ~ 1000 times brighter than a classical nova (and are therefore $10 - 100$ times fainter than a typical supernova).

The second major class of supernovae are the thermonuclear explosions of white dwarfs.

These are known observationally as Type Ia SNe.

The basic physics is related to the Chandrasekhar mass (M_{Ch}): as the white dwarf mass approaches M_{Ch} the central density increases to the point where carbon burning begins. Because the white dwarf is degenerate, the burning is unstable (as in the case of the helium flash). The temperature continues to rise until everything is burnt to nuclear statistical equilibrium (NSE). This all happens faster than the degeneracy pressure can be lifted (unlike the case of the helium flash). The star therefore explodes. Because nuclear reactions generically liberate ~ 1 MeV/baryon, if the entire white dwarf burns this will liberate $\sim 10^{51}$ erg of energy. Most of this energy will be in the form of kinetic energy.

A major outstanding question is the nature of the progenitor system. There are two possibilities: double-degenerate and single-degenerate models. In the former, two white dwarfs merge, pushing the combined mass above M_{Ch} . In the latter, accretion from the secondary companion pushes the white dwarf mass above M_{Ch} . In this case a problem is that accretion onto white dwarfs will tend to produce classical novae. The accreted material burns explosively and is often ejected, thereby failing to actually grow the mass of the white dwarf. As we will see next lecture, there is some evidence for two types of Ia explosions, so perhaps both channels produce explosions.

A second major issue is the precise nature of the burning. This is usually phrased in the literature as deflagration vs. detonation models, where the distinction is in reference to the speed of the burning front (subsonic vs. supersonic). In the case of a subsonic burning front (deflagration), the white dwarf can react (locally) to temperature changes. In particular, the pressure and density will decrease behind the burning front. In the supersonic case, the usual shock jump conditions imply that the density and pressure will increase behind the shock front. The distinction has major implications for the precise nucleosynthesis that occurs as the star explodes. Observational constraints favor a “delayed detonation” model, in which an initial deflagration phase lowers the density somewhat before the detonation occurs. The modeling is, as you can imagine, extremely complicated due to flame physics in degenerate conditions.

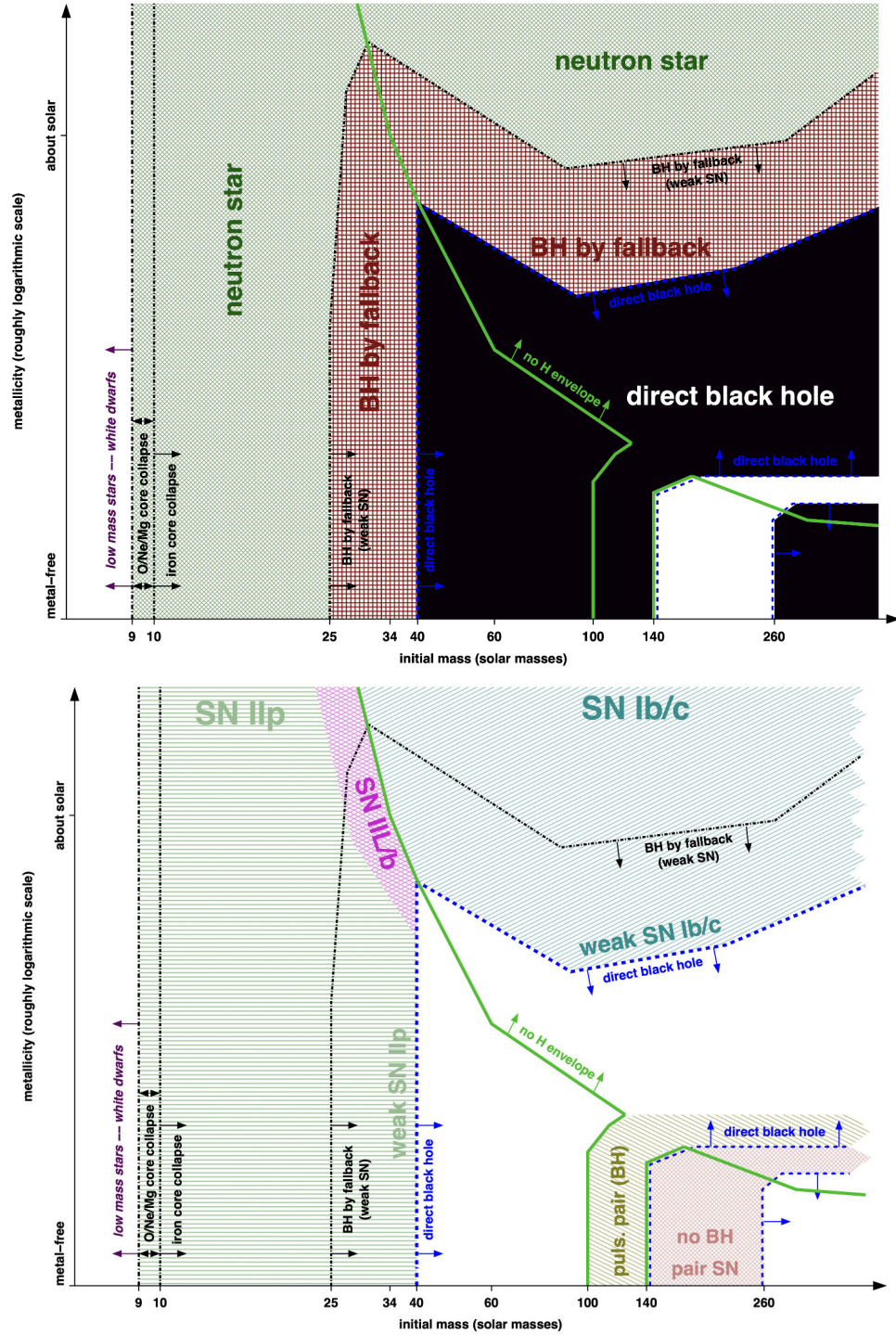


Figure 20.1: Schematic illustration of when compact objects are formed (top panel) and which types of SN they (may) give rise to (bottom panel) as a function of initial mass and metallicity. From Heger et al. (2003). This figure is somewhat out-dated and should only be used to give a sense of the significant variation one might expect as a function of mass and metallicity.

Chapter 21

Supernovae II

In this lecture we will continue our discussion of supernovae, focusing on their observational characteristics including spectral properties, light curves, and demographics.

Supernovae (SNe) are extremely bright transients, with typical peak luminosities of $M_V = -15$ to $M_V = -20$ (with rare super-luminous supernovae reaching peak luminosities of $M_V = -22$) and characteristic timescales of 10 – 100 days (see Fig. 21.1). SNe are brighter than the entire galaxy in which they reside, and can be observed at Gpc distances. Supernovae have been observed since antiquity, including at least eight that were documented in the pre-telescopic historical record. Particularly well-studied examples include the SN of 1054, recorded by the Chinese that produced the Crab Nebula, Tycho's SN in 1572, and Kepler's SN in 1604.

SNe are observationally classified into a variety of types based on optical light curves and spectra. Originally (Minkowski 1941), the classification was based simply on whether or not hydrogen lines were (Type II) or were not (Type I) visible in the SNe spectrum. The modern classification scheme, along with associated progenitors, is as follows:

- **Ia:** No H in spectrum, but prominent Si II absorption. Progenitor: thermonuclear explosion of a white dwarf.
- **Ib:** No H in spectrum, but prominent He absorption. Progenitor: core-collapse SNe (CCSN) from a Wolf-Rayet star lacking a H envelope (could also be due to binarity where mass-transfer results in the stripped envelope).
- **Ic:** No H or He in spectrum. Progenitor: CCSN from a Wolf-Rayet star lacking a

H and He envelope (could also be due to binarity where mass-transfer results in the stripped envelope).

- **IIP**: H in spectrum and a distinct plateau in the light curve. Progenitor: CCSN with a massive envelope.
- **III**: H in spectrum and a linear decline in the light curve. Progenitor: CCSN with a small envelope.
- **IIn**: H in spectrum. Likely CCSN that is interacting with circumstellar material. The “n” indicates narrow emission lines in the spectrum, hence low internal velocity. There are now many “n” type classes, including Ibn, Icn, Idn, Ien.

Fig. 21.1 shows the peak luminosity vs. characteristic timescale for a variety of luminous transients. Several additional (poorly understood) transients not listed above are shown: luminous red novae (LRN), Ca-rich transients, and .Ia explosions. The LRN are believed to arise either (or both) from the merger of main sequence stars, and electron-capture SNe. .Ia explosions are believed to arise from helium detonation on an ultra-compact white dwarf. The parent systems are believed to be AM Canum Venaticorum (AM CVn) systems, in which a degenerate He donor star (e.g., He WD) is transferring mass to a massive CO or ONe white dwarf. The Ca-rich transients are a puzzling category. They are observed far from sites of active star formation, suggesting a WD origin (as opposed to massive stars, which are always found near sites of active star formation). However, they are much less luminous and fade more quickly than Ia SNe, suggesting that they are not the usual detonation of a WD.

Fig. 21.1 also shows the location of **classical novae**. These are luminous eruptions that occur within a binary star system in which a WD accretes material (almost always hydrogen-rich) from a non-degenerate companion. As the accreted layer grows, the density and temperature rise, eventually resulting in ignition of fusion. Much like the thin shell instability we discussed in the context of AGB stars, the burning is a runaway process, with details that depend on the WD mass and accretion rate. The runaway burning causes the accreted envelope to expand and eventually become ejected, leading to a nova. Optical spectra indicate that the ejecta is rich in carbon and oxygen, strongly suggesting that the thermonuclear runaway is able to erode (by mixing) some of the original CO WD. There are 20 – 70 novae per year in our Galaxy with light curve decay times of days to months. Their emission is powered by thermal emission from the hot WD which is in turn powered by the nuclear burning. Classical novae occur in binary systems that are often referred to as cataclysmic variables (CVs), i.e., a main sequence-WD binary in which the secondary star is overflowing

its Roche lobe.

Novae have important implications for the ability of WDs to exceed the Chandrasekhar mass. Models predict that the most massive WDs that stellar evolution can produce are $\approx 1.1M_{\odot}$. Since $M_{\text{Ch}} = 1.4M_{\odot}$, a WD must accumulate $0.3M_{\odot}$ of additional material in order to reach the Chandrasekhar mass and subsequently explode. In the “single degenerate” channel, this is accomplished by accretion from a non-degenerate companion. The key question is whether or not a WD can gently accumulate material. Observations seem to indicate that WDs have a hard time growing via accretion (especially since spectra of novae indicate carbon and oxygen is being removed from the WD during the detonation). But perhaps not all accreting WDs go through a novae phase – the outcome depends strongly on the mass of the WD and the accretion rate of the donor star. This is an active area of research.

There are several distinct physical processes that shape the observed light curve for CCSN. The earliest electromagnetic radiation from a SNe occurs on a timescale of seconds to hours and is associated with the “shock breakout” in which the optical depth to the shock front drops below unity. The timescale for this to occur is roughly R/v_{shock} , implying that we are much more likely to witness a shock breakout from a red supergiant than from a white dwarf. At this point the shock has a temperature of $\sim 100,000$ K, and so the bulk of the radiation comes out at short wavelengths. The rapid expansion and adiabatic cooling of the shock front sets the short timescale. The result is a sudden increase in brightness, first in X-rays and UV, followed by optical emission. Because of the very short timescales involved, only a few shock breakout events have been observed.

The next stage, relevant for CCSN that have exploded within a massive envelope, is governed by recombination radiation from the H-rich envelope. Recombination occurs when the envelope has reached a temperature of ≈ 5500 K (through adiabatic expansion). The temperature and radius at which this occurs is approximately constant, setting the constant luminosity (the plateau). The timescale is set by the time it takes for the entire envelope to recombine (e.g., ~ 100 d).

The linear phase of the light curve is powered by radioactive decay of ^{56}Ni : $^{56}\text{Ni} + e^{-} \rightarrow ^{56}\text{Co} + \nu$ and $^{56}\text{Co} + e^{-} \rightarrow ^{56}\text{Fe} + \nu$. These reactions have half-lives of 6.1 d and 77.1 d, respectively. From $0.07M_{\odot}$ of ^{56}Ni 1.3×10^{49} erg is released. The very late-time light curve (> 1000 d) is powered by the decay of ^{44}Ti . Recent research has explored other power sources beyond radioactive decay, including shocks from interaction with pre-existing mass-loss and energy injection from the newly formed compact object (e.g., a magnetar).

It is clear from the above that the light curve is governed by a variety of parameters: the energy of the explosion, the mass of the envelope, the radius of the pre-supernova star, and the mass of ^{56}Ni produced.

We can put some of this into context by considering the state of the SN immediately after the explosion. At this point, $E \sim 10^{51}$ erg (internal energy), $T \sim 10^{10}$ K, and $R \sim 10^8$ cm. Radiation cannot diffuse out of the star until the age of the system is comparable to the diffusion time, i.e., $\tau_{\text{diff}} \sim t$. From earlier lectures we know that the diffusion timescale is $\tau_{\text{diff}} \sim R^2 \kappa \rho / c$. Setting $\kappa = \kappa_{\text{es}} \approx 0.1 \text{ cm}^2 \text{ g}^{-1}$ for electron scattering (relevant at these high temperatures), $R = vt$, $v \approx 7000 \text{ km s}^{-1}$ (and assuming free expansion so that $v = \text{constant}$), $\rho = \frac{3M}{4\pi R^3}$ and $M = 1M_{\odot}$, we arrive at $t \approx 2 \times 10^6 \text{ s}$ (20 d), and $R = 10^{15}$ cm. In other words, the SNe has expanded by a factor of $\sim 10^7$ by the time that radiation can diffuse out of the expanding fireball. During this phase expansion is adiabatic since energy cannot escape. A photon gas at constant entropy has $T^3/\rho = \text{constant}$, hence $T \propto R^{-1}$. The temperature has therefore also dropped by a factor of $\sim 10^7$ during this time and hence the internal energy is negligible by the time that photons can diffuse out. This implies that radioactive decay is essential to power the light curve.

Detailed models of SN light curves are an active area of research, as they require coupling 3D models of the explosion to sophisticated radiative transfer models and nuclear decay networks.

Type Ia SNe play a critical role in our understanding of the Universe because they are standard candles. Actually, that's not quite right - Ia SNe are *not* standard, but rather *standardizable* candles. As can be seen in Fig. 21.1, Ia SNe do not all have the same peak luminosity, but they do obey a very well-defined relation between peak luminosity and timescale. This is known as the **Phillips relation** (Phillips 1993 – an example of a single-figure, four page paper that has > 1500 citations!). The Phillips relation tells us that brighter Ia SNe have longer characteristic timescales (broader light curves). Therefore, by measuring the light curve, we can deduce the expected luminosity, and therefore the distance.

The physical origin of the Phillips relation is not completely understood. The basic idea is that the quantity varying “along” the relation is the amount of ^{56}Ni mass in the ejecta. More ^{56}Ni will result in a brighter SNe (for reasons outlined above). But what about the timescale? The width of the light curve is likely related to the photon diffusion timescale, which is in turn related to the opacity. Various scenarios for the Phillips relation appeal to higher opacity for larger nickel mass (higher temperature, more iron opacity lines, etc.).

Why does the explosion of such a uniform object ($1.4M_{\odot}$ CO WD) produce such a wide range in ^{56}Ni mass? This is also unknown, but is perhaps related the detailed 3D behavior of the ignition of the WD (e.g., the stochastic and asymmetrical nature of the explosion; see Kasen et al. 2009). In a review article, Hillebrandt & Niemeyer (2000) summarize the situation succinctly: “it appears to be a miracle that all the complexity seems to average out in a mysterious way to make the class so homogeneous.”

In recent years a larger number of “anomalous” Ia’s have been identified, including events much fainter and much brighter than standard Ia’s. The origin of these events is unclear, but is likely a consequence of one or more binary accretion channels.

Very crudely, the energetics of CCSN and Ia SNe are similar in terms of the kinetic energy of the explosion ($\sim 10^{51}$ erg) and the energy in the radiation field ($\sim 10^{49}$ erg). The primary difference is that CCSN also have 10^{53} erg in neutrinos. The similarity in energy scales is largely a coincidence (and very approximate anyway).

The CCSN rate is approximately one per $100M_{\odot}$ formed. This comes from simply integrating the initial mass function (IMF) over the range of masses that we think give rise to CCSN. For example, if a galaxy has a star formation rate of $1M_{\odot} \text{ yr}^{-1}$ (like the Milky Way), then we can expect one CCSN per 100 years. But what initial masses *do* give rise to CCSN? Remarkably, only a few dozen SNe have good *HST* pre-explosion imaging. Such data allow for (a complicated, model-dependent) estimate of the initial mass of the progenitor star by comparing to stellar evolution tracks. There is some debate as to the range, but current data suggest that CCSN are associated with initial masses in the range $8 \lesssim M \lesssim 20M_{\odot}$ (Smartt 2009). More massive stars presumably collapse to form black holes.

The rate of Ia SNe is rather more complicated. In the Galaxy, the rate is ~ 1 per century (how do we know this?). But we know that this rate depends strongly on the age of the underlying stellar population. The Ia SNe rate function is known as the **delay time distribution** (DTD). The DTD quantifies the rate of Ia SNe as a function of the age of the population. Imagine for example a single burst of star formation; the DTD quantifies the rate of Ia’s since the beginning of the burst. Measuring the DTD in observations therefore requires knowledge of the star formation history (SFH) of the sample, as the observed rate, R_{Ia} , is a convolution of the DTD and the SFH:

$$R_{\text{Ia}}(t) \propto \int_0^t \text{SFR}(t-t') \text{DTD}(t') dt' \quad (21.1)$$

Observationally the DTD scales as t^{-1} for ages $\gtrsim 10^8$ yr. The $\sim 10^8$ yr threshold is set by

the time it takes the first white dwarfs to form (i.e., the lifetime of a $\sim 8M_{\odot}$ star).

Where does the t^{-1} dependence come from? If Ia are due to WD-WD mergers (the so-called “double-degenerate” channel), then the merger rate is set by the timescale of gravitational wave inspiral, which scales as $t \sim a^4$ (see the lecture on GWs) where a is the initial semi-major axis of the binary orbit. If the initial separation distribution is $\frac{dN}{da} \propto a^b$ then the merger rate is

$$\frac{dN}{dt} = \frac{dN}{da} \frac{da}{dt} \propto t^{(b-3)/4} \quad (21.2)$$

The observed rate is recovered if $b = -1$. Note that $b = -1$ corresponds to a flat log semi-major axis distribution, i.e., $\frac{dN}{d \ln a} = \text{constant}$. Does nature produce such a distribution? We don’t know! (The distribution is set at least in part by common envelope evolution; hence the uncertainty.) Models have also been proposed to explain the observed DTD distribution within the single-degenerate channel, although they tend to fare less well. It is also possible that the DTD is a composite of two populations, e.g., the single-degenerate channel at early times and the double-degenerate channel at late times. Much is uncertain!

Summary: We don’t really know how CCSNe explode, but when they do they produce a large variety of light curves. This variety is largely a consequence of the envelope mass and hence of the initial mass, metallicity, and mass-loss rates of the progenitor. Given an envelope and the initial conditions of the shock, the subsequent evolution is relatively straightforward. The situation with thermonuclear SNe is much less clear. We are reasonably confident that these are the product of exploding white dwarfs, and empirically most Ia SNe show very well-behaved light curves (obeying the Phillips relation). The progenitors of Ia are still unknown, although there is no shortage of ideas (several channels are likely required to explain the extant data). Their manner of explosion is also uncertain, and the rate at which they explode over time is only weakly empirically constrained.

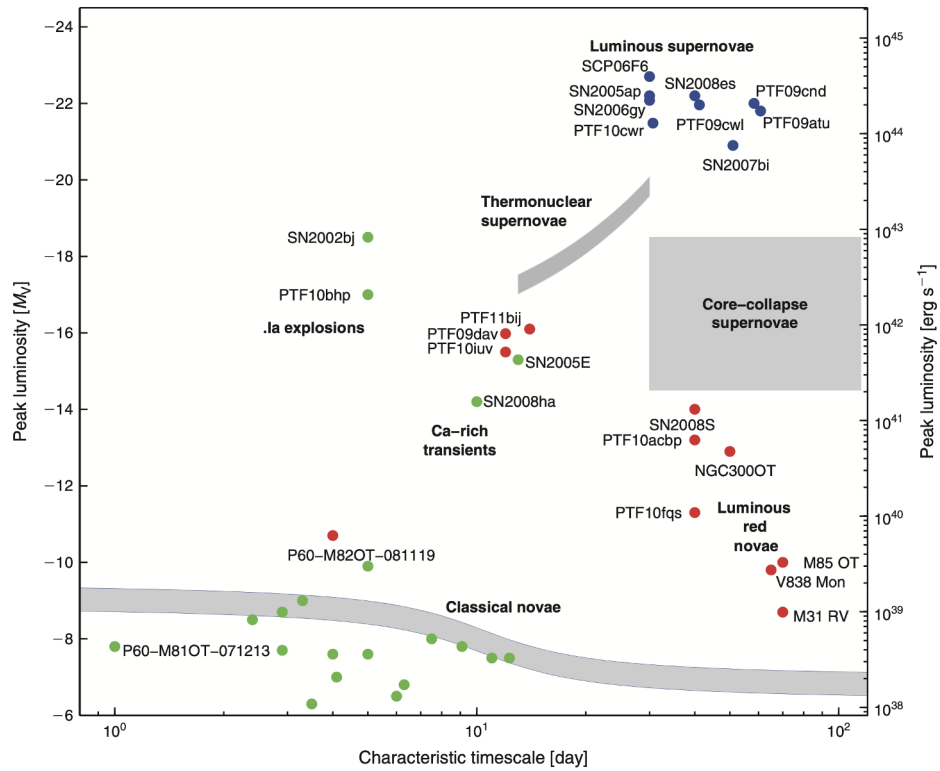


Figure 21.1: Peak luminosity vs. characteristic timescale for a variety of luminous transients. From Kasliwal (2012).

Chapter 22

Gravitational Waves

Gravitational waves (GWs) are waves (propagating fluctuations) of gravitational fields, i.e., “ripples” in spacetime. They emerge as a consequence of General Relativity (GR) and have no Newtonian analog. In order to understand GWs we must therefore appeal to GR throughout this lecture. In some cases results from GR will simply be asserted rather than derived. The derivations can be found in any standard textbook on GR.

The discovery and now routine detection of GWs is a remarkable achievement made possible by LIGO. It is not an exaggeration to state that an entirely new window into the universe has now been opened.

Let's start with the Einstein field equation:

$$R_{\mu\nu} - \frac{1}{2}g_{\mu\nu}R = \frac{8\pi G}{c^4}T_{\mu\nu} \quad (22.1)$$

where $R_{\mu\nu}$ is the Ricci curvature tensor that describes the curvature of spacetime. $T_{\mu\nu}$ is the stress-energy tensor, R is the trace of the Ricci tensor, i.e., $R = \text{tr}(R_{\mu\nu})$, and $g_{\mu\nu}$ is the metric tensor. For example, in flat space-time we have the Minkowski metric:

$$g_{\mu\nu} = \eta_{\mu\nu} = \text{diag}(-1, 1, 1, 1) \quad (22.2)$$

while the Schwarzschild metric is

$$g_{\mu\nu} = \text{diag}\left(-\left(1 - \frac{2GM}{rc^2}\right), \left(1 - \frac{2GM}{rc^2}\right)^{-1}, r^2, r^2 \sin^2\theta\right) \quad (22.3)$$

which describes the gravitational field as a function of distance r outside of a spherical, non-rotating body of mass M . The stress-energy tensor of a perfect fluid in thermodynamic

equilibrium is a diagonal matrix: $T_{\mu\nu} = \text{diag}(\rho, P, P, P)$ where ρ and P are the density and pressure.

Let's consider a small metric perturbation of the form: $g_{\mu\nu} = \eta_{\mu\nu} + h_{\mu\nu}$ where $|h| \ll 1$ and $\eta_{\mu\nu}$ is the Minkowski metric. Plugging this into the field equation leads to a wave equation:

$$\left(\frac{1}{c^2} \frac{\partial^2}{\partial t^2} - \nabla^2 \right) h_{\mu\nu} = -\frac{16\pi G}{c^4} T_{\mu\nu} \quad (22.4)$$

The waves that are the solution to this equation are known as gravitational waves.

By solving the above wave equation (which involves a bunch of GR) one can derive an expression for the *amplitude* of the wave:

$$h \sim \frac{2G}{d} \frac{\ddot{Q}}{c^4} \quad (22.5)$$

where d is the distance from the source, $\ddot{x}$ indicates a second time derivative, and Q is the moment of inertia, i.e., the quadrupole moment of the mass:

$$Q^{ij} = \int \rho(t, \vec{x}) x^i x^j d^3x \quad (22.6)$$

The implication is that there must be time-varying asymmetry in the mass distribution to produce gravitational waves.

Lower-order moments of the mass distribution are not time-dependent. For example, the mass monopole ($\int \rho d^3x$) is time-independent because of mass conservation, and the mass dipole ($\int \rho x d^3x$) is time-independent because of angular momentum conservation.

Let $Q = sMR^2$ for a source of size R and mass M where s represents a dimensionless asymmetry factor ($s = 0$ for symmetric matter; note that an equal mass binary system separated by distance a will have a moment of inertia $Q = 0.5Ma^2$). Then $\ddot{Q} \sim sMv^2$ where v is a characteristic velocity, i.e., $v \sim \dot{R}$ (and setting $\ddot{R} = v^2/R$). Then we have:

$$h \sim \frac{r_s}{d} \left(\frac{v}{c} \right)^2 s \quad (22.7)$$

where $r_s = \frac{2GM}{c^2}$ is the Schwarzschild radius.

What is h , exactly? It is the amplitude of the gravitational wave, which is the stretching and compressing of spacetime. If we imagine two objects separated by distance x then h determines the *fractional* change in the distance (dx/x) between these two objects as the wave passes.

Consider a pair of merging neutron stars. In this case $r_s \sim 1$ km. Let $v \sim c$. The distance to the event is d , and a typical value for this will depend on the very uncertain event rate (if

merging neutron stars are common then we can expect to detect them at closer distances for a fixed observing period). Models predict a rate in the ballpark of ~ 1 per year per several hundred Mpc^3 . This would suggest a value for d of $\gtrsim 100 \text{ Mpc} \sim 10^{26} \text{ cm}$ if we hope to observe one or more per year. This implies a typical value for h of $r_s/d \sim 10^5 10^{-26} \sim 10^{-21}$. This is known as the dimensionless **strain** and is the quantity that is directly observable.

Notice that $h \sim d^{-1}$, so the signal does not decrease with distance as fast as for the signal from photons. This is because h measures the amplitude of the wave while h^2 measures the energy (as is usual with waves of all kinds). With EM waves we are usually sensitive to the energy in the field, which scales as d^{-2} . But for GWs we can directly measure the amplitude. This weaker scaling with distance than we are used to has important implications. For example, if we can increase the sensitivity of our detectors by $2\times$ then we can probe $8\times$ the volume and hence should expect $8\times$ more events. That is a huge gain! Compare this with EM waves, in which the d^{-2} scaling means that $2\times$ greater sensitivity results in an increase of only $2^{3/2} \approx 2.8$ in the volume and hence the rates.

As we have seen, a typical value of the strain is $h \sim 10^{-21}$. How small is this? Imagine a measurement device of length 1 km. In order to detect $h \sim 10^{-21}$ we would need to measure a distance shift of 10^{-16} cm . Recall that the size of the nucleus is $\sim 1 \text{ fm} = 10^{-13} \text{ cm}$. So we would need to measure a distance change that is 10^{-3} the size of the nucleus! Obviously this is not an easy measurement.

There are two primary means of measuring the strain. The first is via interferometry. LIGO is a Fabry-Perot interferometer using a $1\mu\text{m}$ laser in a cavity of length 4 km. The light beam bounces within the cavity approximately 100 times before exiting due to the high reflectivity of the mirrors. Assuming $h = 10^{-21}$, the phase shift of the light in this case is:

$$\Delta\phi = 100\Delta L \frac{2\pi}{\lambda} = 100 \times 4 \times 10^{-16} \text{ cm} \frac{2\pi}{1\mu\text{m}} = 2.4 \times 10^{-9} \quad (22.8)$$

which is a very small shift. However, it can be measured if the photon Poisson (shot) noise is below this level. Recall that Poisson uncertainty on the mean scales as $1/\sqrt{N}$, so we require $\sim 10^{18}$ photons over a time interval of 0.01 s (for a 100 Hz gravitational wave). This corresponds to a 100 W laser (LIGO uses a 200 W laser). Higher power lasers directly translates to higher SNR GW measurements.

Many sources of noise exist in the instrument. Some can be averaged away, but others are environmental (e.g., seismic signals). This is one of the reasons for two separate observatories (another reason is better sky localization; see below).

The second main detection technique is via pulsar timing arrays. A very low frequency GW

will induce a wobble in the location of the Earth relative to distant stars. This wobble will induce a change in the arrival time of pulses from pulsars. In 2023 the first detection of GWs via this technique was announced. The result was a detection of the GW background – a statistical result based on the sum total of many faint GW signals.

Let's consider in greater detail the GW signal from two orbiting compact objects (e.g., NS-NS, NS-BH, or BH-BH). Imagine the pair is orbiting with frequency $f = 2\omega_{\text{orb}}$, with semi-major axis a and reduced mass $\mu = m_1 m_2 / (m_1 + m_2)$. The amplitude of the GW is then:

$$h_0 \sim \frac{2G\mu\omega_{\text{orb}}^2 a^2}{dc^4} = \frac{2G^2 m_1 m_2}{da c^4} \quad (22.9)$$

where $h(t) = h_0 \sin(ft + \phi)$ and the second equality comes from invoking Kepler's third law.

From GR we learn that the luminosity in GWs is:

$$L_{\text{GW}} = -\frac{dE}{dt} = \frac{1}{5} \frac{G}{c^5} \langle \ddot{Q}_{ij} \ddot{Q}^{ij} \rangle \quad (22.10)$$

where $\langle \cdot \rangle$ denotes an average over a time longer than the characteristic period of the source.

In terms of the binary system this can be expressed as:

$$L_{\text{GW}} = \frac{32}{5} \frac{G^4}{c^5} \frac{1}{a^5} m_1^2 m_2^2 (m_1 + m_2) \quad (22.11)$$

which is roughly $L_{\text{GW}} \sim G^4 c^{-5} a^{-5} M^5$. If we set $a = r_s$ then $a \sim GM/c^2$ and so $L_{\text{GW}} \sim c^5/G$ which is $\sim 10^{59} \text{ erg s}^{-1}$. That is a lot of luminosity.

By energy conservation, the energy lost via GW emission must come from somewhere. In the case of a binary system it comes at the expense of the orbital energy: $E_{\text{orb}} = -\frac{Gm_1 m_2}{2a}$.

If we set $L_{\text{GW}} = -dE_{\text{orb}}/dt$ then we have:

$$\frac{da}{dt} = -\frac{64}{5} \frac{G^3}{c^5} a^{-3} m_1 m_2 (m_1 + m_2) \quad (22.12)$$

From this expression we can derive a timescale for the orbital separation to shrink to zero, i.e., for the binary to merge as a consequence of GW emission:

$$t_{\text{GW}} = \int_0^{t_{\text{GW}}} dt = \int_{a_0}^0 \frac{dt}{da} da = \frac{5}{256} \frac{c^5}{G^3} \frac{a_0^4}{m_1 m_2 (m_1 + m_2)} \quad (22.13)$$

This is a long timescale for ordinary length and mass scales. For example, for $m_1 = m_2 = M_{\odot}$ and $a_0 = 1 \text{ AU}$, $t_{\text{GW}} = 2 \times 10^{17} \text{ yr}$. But the timescale is a very strong function of a , so if binaries can be brought much closer together by another process, e.g., dynamical interactions or common envelope evolution, then GW emission can result in a merger within a Hubble

time (10^{10} yr). This is one of the reasons that common envelope evolution is such an important physical process to understand.

The decay rate predicted by Equation 22.12 was measured for the first time in 1974 in a pulsar-neutron star binary by Hulse and Taylor. This was the first indirect evidence for the existence of GWs and for their work they were awarded the Nobel Prize in 1993.

We can estimate the frequency, f , of the GW at (or near) the time of the merger by considering the last stable circular orbit around a compact object. For a BH this is the inner-most stable circular orbit (ISCO), which for a non-rotating BH is $r_{\text{isco}} = 3r_s$, so:

$$f_{\text{GW}} = 2\sqrt{\frac{G(m_1 + m_2)}{r_{\text{isco}}^3}} = 460 \text{ Hz} \frac{60M_{\odot}}{m_1 + m_2} \quad (22.14)$$

More generally, the following scaling relations hold throughout the inspiral and merger: $f \propto a^{-3/2}$, $h \propto a^{-1}$ and $a = (a_0^4 - Ct)^{1/4}$ where C is a constant. The early evolution of the inspiral occurs at lower frequency. The maximum frequency occurs for only a short period of time (less than a second for stellar mass compact objects).

The full behavior of the GW emission is separated into three phases: inspiral, merger, and ringdown. The first and last of these can be reasonably well-described by analytic post-Newtonian and GR calculations. The merger phase is complicated and requires numerical GR simulations. LIGO detects the merger and ringdown phases, so the large-scale computation of numerical relativity calculations has been an essential ingredient in interpretation of the data (a bank of “filters” based on such calculations are used to analyze the data). The earlier inspiral phase occurs at a frequency that is too low for LIGO to detect.

In principle, the measurement of $h(t)$, which LIGO provides, allows the determination of m_1 , m_2 , a and d as well as other quantities such as the spin of the BH.

The observational landscape for detecting GWs is rich and rapidly changing. LIGO is a joint project between MIT and Caltech funded by NSF in the 1990s with sites in Washington and Louisiana. Data collection occurred between 2002 and 2010 and resulted in zero detections. Advanced-LIGO was a substantially upgraded system that was completed in 2015. The third observing run (O3) was just recently completed (in late 2020). aLIGO is now routinely detecting compact binary mergers. VIRGO is a GW observatory in Italy that came online after LIGO and is currently less sensitive than LIGO, but is useful for localization of sources. KAGRA is a Japanese interferometer that has recently been commissioned. Increasing the number of observatories substantially increases the localization precision of the source on-

sky.

Pulsar timing arrays (PTAs) use the fact that pulsars are extremely stable clocks. A large network of pulsars can be used to detect passing gravitational waves, as the passing waves will cause slight changes in the arrival time of the pulses. PTAs are less sensitive than LIGO, but are uniquely sensitive to very low frequency waves ($< 10^{-6}$ Hz). The PTA is sensitive to the stochastic GW background composed of the superposition of many waves from the inspiral phase of supermassive black hole binaries at the centers of galaxies.

LISA is a planned ESA-led mission to detect GWs via a space-based interferometer. It will have much longer baselines and will be sensitive to much lower frequencies and hence much larger mass BH mergers (SMBHs).

The detection of signals in the LIGO data is very complex. The very first discovery (GW150914) was much higher SNR than expected, and so the signal was relatively easy to see directly in the data. Most sources are lower SNR, and so signal processing is required. The basic idea is that a library of templates based on numerical relativity calculations are used to filter the data. This is known as “matched filtering”.

Source localization is a major problem in GW astronomy. A single interferometer is sensitive to waves across the entire sky (very much unlike conventional telescopes which must point to a particular patch of the sky). Note that the wavelength of a typical 100 Hz gravitational wave is 3000 km! This all-sky nature of GWs is both a blessing and a curse. One can simultaneously view the entire sky, which is good for identifying events, but it makes localization impossible with a single interferometer. With more than one detector one can use the finite time travel of the wave to improve localization. Even with three separate observatories the sky localization is not better than $\sim 10 - 100$ sq. deg. There is substantial interest in detecting the electromagnetic counterpart to the GW event (if such a counterpart exists), which demands rapid followup and monitoring of significant regions of sky.

GWs cannot be used to directly image the source, since the wavelength of the wave is comparable to the size of the source, unlike EM waves, which have wavelengths much smaller than the source size.