

Cosmology & Galaxy Formation

Charlie Conroy

v1, September 6, 2026

Preface

This book forms the core of AY202a, a graduate course I have taught at Harvard since 2015. Principal sources for this material include “Introduction to Cosmology” by Ryden, “Galactic Dynamics” by Binney & Tremaine, “Galaxy Formation and Evolution” by Mo, van den Bosch, and White, “Galaxy Formation” by Longair, “Nucleosynthesis and Chemical Evolution of Galaxies” by Pagel, and lecture notes from Daniel Eisenstein, David Weinberg, and Mark Krumholz. I thank the former students of AY202a for providing feedback and identifying errors in previous versions.

The contents are organized as follows. Part I provides a broad overview of concepts, basic observational facts, and the importance of timescales. Part II discusses the dynamics of the homogeneous universe, i.e., the big bang and inflation, cosmic acceleration, and the cosmic microwave background. Part III is a detour into the fundamentals of gravitational dynamics. This material is essential for understanding the topics in the latter part of the book. Part IV describes the growth of structure, from linear perturbations to non-linear evolution and the formation of gravitationally-bound structures. Finally, Part V addresses the topic of galaxy formation and evolution. This book assumes basic undergraduate-level knowledge of astronomy and physics.

Contents

I	Setting the Stage	1
1	Introduction	2
II	Homogeneous Cosmology	12
2	Cosmological Dynamics	13
3	Cosmological Tests	24
4	The Cosmic Microwave Background	33
5	Cosmic Inflation	41
III	Gravitational Dynamics	47
6	Potential Theory & Relaxation	48
7	Distribution Functions and the Jeans Equations	58
8	The Virial Theorem & Dynamical Friction	66
IV	Growth of Structure	76
9	Linear Fluctuations	77
10	Time and Scale-Dependence	84
11	Non-Linear Evolution & Spherical Collapse	93
12	Dark Matter Halos	101
13	Dark Matter	109

V Galaxy Formation & Evolution	119
14 Galaxies: Key Observational Facts	120
15 Gas and the Road to Star Formation	127
16 Star Formation	139
17 Stellar Populations	150
18 Chemical Evolution	157
19 Stellar & Black Hole Feedback	169
20 Statistics of the Galaxy Distribution	178
21 Successes and Challenges of Current Models	188

Part I

Setting the Stage

Chapter 1

Introduction

A quantitative understanding of the evolution of our Universe from the Big Bang to the emergence of stars and galaxies is one of the hallmark achievements of 20th century science. The story involves many threads, including cosmological models emerging from general relativity, the inference of a hot Big Bang, the detection and interpretation of the cosmic microwave background (CMB), the development of tiny density fluctuations and their subsequent evolution under gravity, the cooling and heating of baryons, and the formation and evolution of galaxies.

This course is a survey spanning all of these topics (and more). We will not be able to cover every topic in the broad fields of galaxy formation and cosmology, nor any particular topic in great depth. The goal is to provide, in one course, a relatively complete view of the story of the universe.

The course is titled “Extragalactic Astronomy & Cosmology” – what does that mean? **Cosmology** is the study of the content and history of the Universe averaged over large temporal and spatial scales, e.g., ~ 100 Mpc. **Extragalactic Astronomy** refers to the study of phenomena outside of the Milky Way, i.e., the formation and evolution of galaxies. In other words, the course covers everything in the Universe outside of our Galaxy, although we will have occasion to comment on processes within our Galaxy as well.

One way to frame the goals of this field is that we are attempting to solve an initial value problem. In principle the solution is straightforward, as we have a set of equations (all of physics) and an initial condition (the Big Bang) and our goal is to ‘solve’ this problem in order to understand the evolution of the Universe and the material within it. However, we do not know the initial conditions all that well, as that requires understanding of the Big Bang, inflation, and probably a unification of general relativity and quantum mechanics. Moreover,

while we generally know the equations relevant for the subsequent evolution (general relativity, quantum mechanics, the laws of thermodynamics, etc.) there is no hope of directly solving this coupled set of equations for all particles in the universe. To make progress, one either breaks off a self-contained subset of the problem or tries to solve a simpler set of equations that approximates the full physics. How this is done in practice is of course a major component of modern research.

It is worth remembering that not all problems in physics are initial value problems. For example in thermodynamics, or in general when the system has had time to *relax* (we will define this term in a later chapter), the system loses all memory of its initial conditions and converges to simple distributions, e.g., a Maxwellian in the case of a thermal distribution of particles. In most cases that we will encounter, the initial conditions (ICs) are imprinted on the final state. This has advantages, as it means we can learn about the ICs from present-day observations.

Galaxies are the largest discrete objects in the Universe and are spectacular examples of emergent phenomena, in which complex physical processes and laws give rise to objects that are remarkably regular and uniform (see Fig. 1.1). They reflect the interaction of many physical processes including: gravity, mergers/interactions, gas cooling, star formation, feedback from stars and black holes, radiation pressure and ionization, cosmic rays. Each of these processes is comparably important. The field of galaxy formation thus lies at the intersection of many disciplines; this makes the field both very challenging and very stimulating.

The observational picture is very incomplete. Only relatively recently has panchromatic data become available, and even then mostly for objects at low redshift. Only for galaxies in the local group (and slightly beyond) can we resolve individual stars; for all other systems we must interpret the combined light from all the stars in the galaxy, an approach called stellar population synthesis. From an observational perspective, “galaxy evolution” is a story of snapshots observed at different redshifts that are difficult to connect together (compare and contrast with archeology). Dark matter is essential for connecting the galaxies across time, but is very difficult to even indirectly observe (and of course we still have no idea what dark matter *is*). Many of the most important physical processes, such as stellar and black hole feedback, are almost entirely hidden from us.

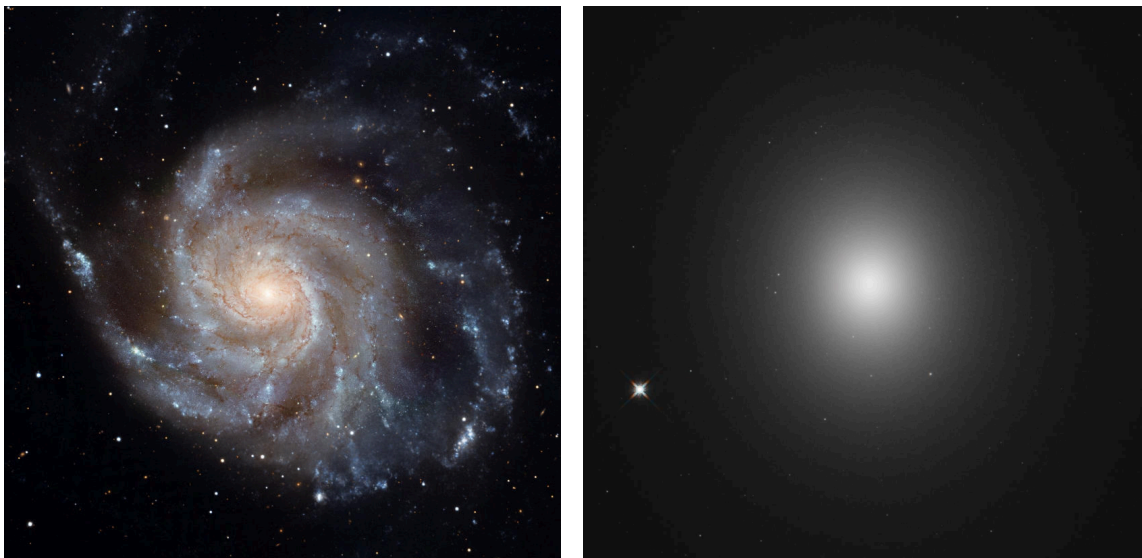


Figure 1.1: M101 (left) and M49 (right) representing a typical grand-design spiral galaxy and an elliptical galaxy. Images courtesy of ESA/Hubble.

Resolving Galaxies: Galaxies have sizes of $\sim 1 - 10$ kpc (their light profiles are exponential, so by “size” we mean a characteristic exponential scale). For very nearby systems where cosmological effects can be ignored, one arcsecond ($1''$) corresponds to 5 pc at a distance of 1 Mpc and 50 pc at 10 Mpc. At $z > 1$ this scale asymptotes to ≈ 8 kpc/ $''$. Since the typical seeing of ground-based optical observations is $0.5 - 1''$, this means that high-redshift galaxies are only marginally resolved from the ground. Observations with *Hubble Space Telescope* (*HST*) and the *James Webb Space Telescope* (*JWST*), which have angular resolution of $\approx 0.1''$, are therefore essential to understand the internal structure of galaxies.

Before embarking on our semester-long endeavor, we would do well to have some sense of our goals. In extragalactic astronomy the basic goals are twofold: *to discover phenomena (observations) and to understand the phenomena in terms of physical laws (theory)*. In cosmology, the field has mostly focused on the goal of measuring cosmological parameters and testing for deviations from the concordance cosmological model (which we will define in later chapters). An important question that arises in any intellectual project, including the present one, is this: *how good is good enough?* How precisely do we want to understand the story of the universe before we declare victory and move on? It’s worth keeping this question in mind throughout the course.

There are a set of fundamental observational facts that ground our understanding of the universe:

1. The night sky is dark. If we assume the universe is infinite and isotropic, and uniformly populated by stars with average luminosity L and number density n , then the sky brightness due to a shell of stars of thickness dr is:

$$dJ = \frac{L}{4\pi r^2} nr^2 dr = \frac{Ln}{4\pi} dr \quad (1.1)$$

where J is the intensity of radiation, i.e., energy per time per unit area per steradian on the sky. The total brightness of the sky can be computed by integrating over all shells:

$$J = \frac{Ln}{4\pi} \int_0^\infty dr \rightarrow \infty \quad (1.2)$$

This is known as Olbers' Paradox. It is a paradox because the night sky is dark, not infinitely bright! Possible solutions to this paradox include:

- (a) the universe is not infinite (the integral does not extend to infinity).
- (b) dust absorbs/obscures the light (this does not work, as eventually the dust would heat up to the temperature of the starlight).
- (c) the universe is not infinitely old, which, when combined with the finite speed of light, implies an effective upper limit to the integral.
- (d) flux does not decline as r^{-2} .

It turns out that (c) is the primary solution, although (d) plays a nontrivial role due to redshifting effects.

2. The universe is homogeneous and isotropic (on sufficiently large scales). This is the **cosmological principle**. The CMB is isotropic to a level of $\sim 10^{-5}$. It is very contrived to imagine a universe that is isotropic but not homogeneous (e.g., if we were at the top of a cosmic mountain, the universe would be isotropic but not homogeneous). All observations to-date agree with the cosmological principle.
3. The universe is expanding. The Doppler shift allows us to relate wavelength shifts to velocities via:

$$z \equiv \frac{\lambda_{\text{obs}} - \lambda_{\text{em}}}{\lambda_{\text{em}}} \approx \frac{v}{c} \quad (1.3)$$

where z is the redshift, v is the line of sight velocity, c is the speed of light, and λ_{obs} and λ_{em} are observed and emitted wavelengths. The relation to velocity above holds only for low velocities; as velocities approach the speed of light, one must use the relevant relativistic expression. Vesto Slipher measured redshifts for a sample of nearby galaxies and Edwin Hubble measured distances based on Cepheid standard candles and in 1929 deduced the famous relation between velocities and distances: $v = H_0 r$ where H_0 is known as the Hubble constant. Hubble initially deduced $H_0 \approx 500 \text{ km s}^{-1} \text{ Mpc}^{-1}$, which was incorrect by a factor of ≈ 10 due to the fact that Hubble used two kinds of Cepheid variables and also misclassified some bright stars. The current value of the Hubble constant is $H_0 = 70 \pm 2 \text{ km s}^{-1} \text{ Mpc}^{-1}$ (see Fig. 1.2). Note that the Hubble *constant* refers to the expansion as measured today. As we will see later, the Hubble parameter is a function of time (redshift).

An expanding universe naturally produces a linear distance-velocity relation, and so Hubble’s result strongly ruled out the prevailing paradigm of a static universe. An expanding universe naturally implies both a finite age and extent of the universe, which, on dimensional grounds can be estimated via: $t_H \equiv H_0^{-1} \approx 14 \text{ Gyr}$; $D_H = c/H_0 \approx 4.3 \text{ Gpc}$.

4. The universe is filled with a uniform cosmic background radiation that peaks in the microwave. In 1965 the CMB was serendipitously detected by Arno Penzias and Bob Wilson. It is a nearly perfect blackbody with temperature 2.73 K. An expanding universe implies a relation between radiation temperature and redshift as $T \propto (1+z)$, which implies the CMB was *hotter* in the past. The discovery of the CMB and an expanding universe naturally suggests a universe that began with a **hot Big Bang**.
5. The universe is comprised mostly of H and He, with trace amounts of other light elements such as D, Li and Be. Observations of “primordial” abundances of these light elements are in excellent agreement with standard big bang nucleosynthesis (BBN). This implies that the universe during the first few minutes after the Big Bang is relatively well-understood.

Table 1.1 collects the major events in the history of the universe, several of which we will return to in later chapters.

The evolution of objects in the universe is governed by many disparate **timescales**. Several of the most important for extragalactic astronomy and cosmology are listed below.

Table 1.1. Important events in the history of the universe

Time	Redshift	Event
0.0		Big Bang
10^{-43} s		Planck time; end of quantum gravity (?)
10^{-34} s		inflation ends (?)
10^{-11} s		baryogenesis (?)
10^{-6} s		quarks $\rightarrow$ hadrons ($kT \approx$ nucleon binding energy)
1 s		neutrinos decouple: $t_{\text{univ}} = (\sigma_{\nu} n c)^{-1}$
few s		$e^+ + e^-$ annihilate; heat added to the universe
1 m		light nuclei form; $kT \sim 0.1$ MeV ($\frac{1}{20}$ binding energy of D)
10 m		free neutron half-life
10^3 yr		matter-dominated era begins
10^5 yr	$\sim 1,000$	recombination. Photons decouple from matter. Formation of CMB.
10^7 yr	~ 100	first stars (?)
10^8 yr	$\sim 10 - 20$	first galaxies (?)
10^9 yr	$6 - 10$	reionization
10^{10} yr	≈ 0.5	Λ -dominated (dark energy) era begins

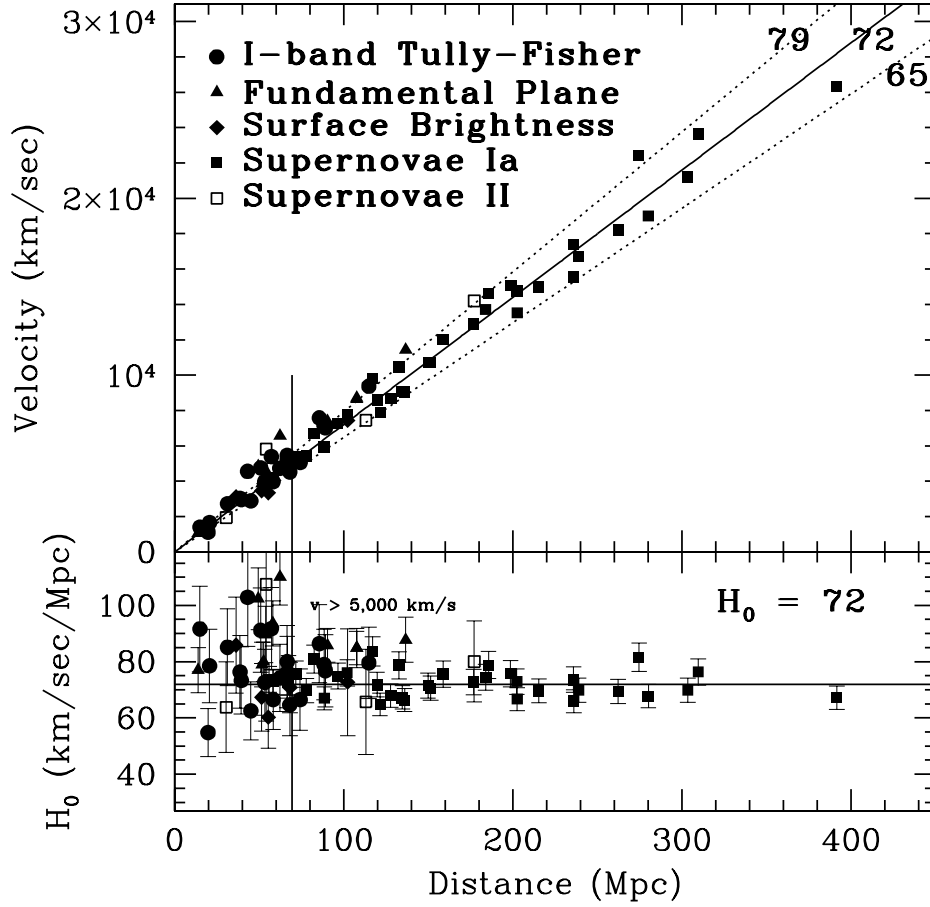


Figure 1.2: State-of-the-art estimate of the Hubble constant circa 2001, based on data from the *Hubble Space Telescope* Key Project (Freedman et al. 2001) as well as more distant distance anchors.

- **Hubble time (t_H):** This is the timescale on which the universe as a whole evolves. It is (up to a constant of order unity) the age of the universe at the redshift of interest.
- **Dynamical time (t_{dyn}):** This is the time required to orbit a dynamical system in equilibrium: $t_{\text{dyn}} = \sqrt{\frac{3\pi}{16G\bar{\rho}}}$ where $\bar{\rho}$ is the mean density of the system. The dynamical time is comparable to the crossing time: $t_{\text{cross}} = r/v$ and the free-fall time: $t_{\text{ff}} = t_{\text{dyn}}/\sqrt{2}$.
- **Cooling time (t_{cool}):** This is the time it takes for a system to radiate away its thermal energy: $t_{\text{cool}} = E/\dot{E}$ where E is the thermal energy and $\dot{E}$ is the energy loss rate, e.g., the luminosity. In some contexts this is known as the thermal timescale.

- **Star formation time (t_{SF}):** This is the time it takes for a system to convert all of its gas into stars: $t_{\text{SF}} = M_{\text{gas}}/\text{SFR}$ where SFR is the star formation rate and M_{gas} is the gas supply available for star formation (usually restricted to the cold gas supply).
- **Dynamical friction time (t_{DF}):** This is the timescale on which a satellite loses orbital energy and sinks to the center and merges with the host system. $t_{\text{DF}} \propto t_H \frac{M_h}{M_s}$ where M_h and M_s are the masses of the host and satellite.

We can make progress in understanding how a system will evolve by considering the hierarchy of timescales. Some examples:

- Processes with timescales longer than t_H are unimportant. For example, $t_{\text{DF}} > t_H$ in massive clusters, hence satellites in clusters do not merge, and clusters therefore contain many satellites. If $t_{\text{cool}} > t_H$ then the gas will not cool and so will be unavailable to form stars.
- When $t_{\text{cool}} > t_{\text{dyn}}$ the gas will settle into hydrostatic equilibrium. An example of this regime is the hot gas in galaxy clusters.
- When $t_{\text{cool}} < t_{\text{dyn}}$ the gas cannot maintain hydrostatic equilibrium and will cool and collapse on a dynamical (free-fall) time.
- When $t_{\text{SF}} \sim t_{\text{dyn}}$ gas will turn into stars during the initial collapse of the system (thought at one time to be the formation channel of early-type galaxies known as “monolithic collapse”).
- When $t_{\text{SF}} > t_{\text{dyn}}$ gas will settle into a rotationally-supported disk and will form stars “slowly”. This is the regime of disk-dominated star forming galaxies like the Milky Way.

Magnitudes are a peculiar astronomical unit that has its origins in early attempts to place the naked-eye stars on a “natural” unit system. The scale is inverted, in the sense that brighter stars have smaller magnitudes than fainter stars, and because our eyes are approximately logarithmic detectors, the resulting scale is approximately logarithmic. Specifically, the magnitude of an object of flux f is: $m \propto -2.5 \log f$ (where the log is assumed base-10). The coefficient is not special or important – it was determined in an attempt to place the historical magnitude scale on a quantitative footing. Notice that a difference in one magnitude corresponds to a 2.5 difference in flux, and a 5 magnitude difference corresponds to a

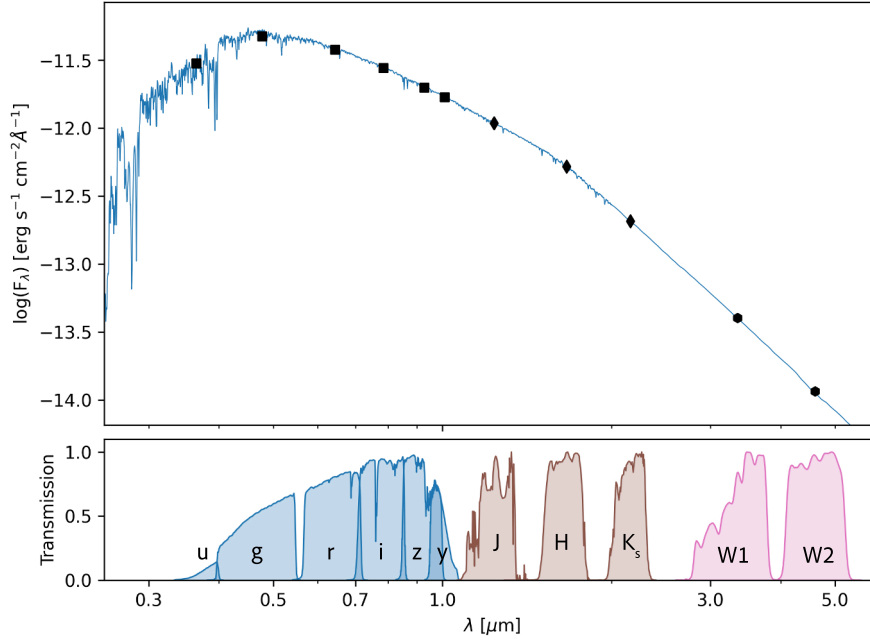


Figure 1.3: Synthetic solar spectrum and corresponding synthetic photometry (points). The bottom panel shows the filter curves used to generate the photometry. The transmission includes the effect of the Earth’s atmosphere for ground-based observatories. Systems include Rubin Observatory (ugrizy), 2MASS (JHK_s) and *WISE* (W1, W2). Figure courtesy of Phill Cargile.

factor of 100 in flux. Notice further that $2.5 \log$ is close to the natural logarithm ($\ln$), i.e., $m \sim -\ln f$.

The above equation requires a constant, also known as a zero-point, i.e., $m = -2.5 \log f + \text{ZP}$. There are two common conventions in use today: 1) use the star Vega to define the zero-points such that $m \equiv 0$ at all wavelengths for the star Vega (an A0V star); 2) the AB system, in which $m_\nu = -2.5 \log f_\nu - 48.60$ (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. Be sure to understand the zero-point system of a particular dataset before using it.

The equation above holds both in a monochromatic sense, e.g., $m_\nu \propto -2.5 \log f_\nu$ and also if one takes observations through a bandpass filter. For simplicity, assume that a bandpass filter is box-shaped with a central wavelength, λ_{eff} and a width $\Delta\lambda$. It is common to define the resolution, R , as $R \equiv \lambda_{\text{eff}}/\Delta\lambda$. Most astronomical filters have $R \sim 5$ (note that

spectrographs have $R \sim 10^3 - 10^5$). There are many filter systems in use, including Johnson-Cousins (UBVRI), SDSS (ugriz), LSST/Rubin (ugrizy), and NIR systems (JHKLM). There are hundreds of filters spanning the UV-IR. When quoting the apparent magnitude in some filter, e.g., the B -band filter, we write $B = 5$ (for example), while when quoting absolute magnitudes we write $M_B = 5$. You are encouraged to take some time to familiarize yourself with the most common filter systems. See Bessell (2005) for an excellent review, and Fig. 1.3 for some examples.

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

$$m_1 - m_2 = -2.5 \log(f_1/f_2) = 5 \log(d_1/d_2) \quad (1.4)$$

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

$$m - M = 5 \log(d/10 \text{ pc}) + A + K \quad (1.5)$$

where $m - M$ is known as the distance modulus (sometimes written as μ), A is the extinction (in magnitudes) between the source and the observer, and K is the K -correction, which corrects for the fact that a fixed observed-frame filter will sample different restframe wavelengths for sources at different redshifts.

Finally, you will sometimes come across discussion of bolometric corrections, which are defined as $BC_X = m_{\text{bol}} - X$ for some filter X . The bolometric correction is the correction required to translate an observation in some filter to the bolometric flux. BCs are a function of the spectral energy distribution (SED) of the source, and so are model-dependent.

Part II

Homogeneous Cosmology

Chapter 2

Cosmological Dynamics

In this chapter we will assemble the most important equations governing the evolution of the homogeneous universe, including the Robertson-Walker metric of spacetime in an expanding universe, and the three equations governing the dynamics of an expanding universe. These latter equations provide us with connections between the composition of the universe (e.g., radiation, matter, dark energy), the curvature of the universe (open, flat, closed), and the expansion history of the universe. In the following chapter we will explore these connections in detail.

The observation that the universe is expanding according to the cosmological principle (isotropically and homogeneously) leads us to define two distinct coordinate systems. Proper, or physical coordinates is the familiar coordinate system in which one measures distances between objects with rulers. The distance between objects expanding with the Hubble flow increases with time in proper coordinates. A second system, known as **comoving coordinates**, is one in which the expansion of the universe is removed (or divided out). In this system, the distance between objects expanding with the Hubble flow does *not* change with time. This can be a confusing coordinate system to work with initially, but we will see its value in the course of this and the following chapters.

The comoving system relies on a concept called the **scale factor**, $a(t)$. Consider uniform expansion: $r(t) = r_0 a(t)$. In this equation, $r(t)$ is the physical coordinate, r_0 is the comoving coordinate, and $a(t)$ is the time-dependent (and spatially-independent, as required by the cosmological principle) scale factor. We can connect the scale factor to the Hubble parameter

by considering velocities: $v = \dot{r} = \dot{a} r_0 = \frac{\dot{a}}{a} r$. Comparison to Hubble's law reveals:

$$\boxed{H(a) = \dot{a}/a} \quad (2.1)$$

As we will see, understanding the evolution of the universe boils down to determining the time-dependence of the scale factor.

We can relate the scale factor to redshift by combining the Doppler formula and Hubble's law:

$$dv = \frac{\dot{a}}{a} dr \quad (2.2)$$

$$\frac{d\lambda}{\lambda} = \frac{dv}{c} = \frac{\dot{a}}{a} \frac{dr}{c} = \frac{\dot{a}}{a} dt = \frac{da}{a} \quad (2.3)$$

which implies $\lambda \propto a$. We can use this result to write: $\frac{\lambda_{\text{obs}}}{\lambda_{\text{em}}} = \frac{a_0}{a_e} = 1 + z$. We have the freedom to define $a_0 \equiv 1$; hence $\boxed{a = (1 + z)^{-1}}$.

A **metric** specifies the distance between two points. You are all familiar with the standard metrics of flat 2D space in Cartesian coordinates:

$$dl^2 = dx^2 + dy^2 \quad (2.4)$$

and polar coordinates:

$$dl^2 = dr^2 + r^2 d\theta^2 \quad (2.5)$$

The following is the metric for the distance between two points on the surface of a sphere of radius R_c :

$$dl^2 = R_c^2 d\theta^2 + R_c^2 \sin^2\theta d\phi^2 \quad (2.6)$$

if we define $r \equiv R_c\theta$ then

$$dl^2 = dr^2 + R_c^2 \sin^2(r/R_c) d\phi^2 \quad (2.7)$$

(compare this to the flat polar case above). In the case of hyperbolic geometry (saddles) one replaces $\sin \rightarrow \sinh$ in the above equations.

In special relativity we have the Minkowski metric for flat spacetime:

$$ds^2 = -c^2 dt^2 + dr^2 + r^2 d\Omega^2 \quad (2.8)$$

where $ds^2 = 0$ for photons and $d\Omega^2 = d\theta^2 + \sin^2\theta d\phi^2$.

The metric for a spatially homogeneous isotropic universe that includes expansion/contraction is given by the **Robertson-Walker metric**:

$$ds^2 = -c^2 dt^2 + a(t)^2 [dr^2 + S_k^2(r) d\Omega^2] \quad (2.9)$$

where r is in comoving units. The term $S_k(r)$ captures the global curvature of space and is equal to $R_c \sin(r/R_c)$ for spherical geometry (positive curvature), r for flat geometry, and $R_c \sinh(r/R_c)$ for hyperbolic geometry (negative curvature). In the first and last cases, R_c is the radius of curvature. If you had taken this course 20 years ago a lot of time would have been spent deriving equations assuming different global geometries. Thankfully, the CMB tells us that the universe is flat to extremely high precision so we will assume flat geometry in most of the course.

The Robertson-Walker metric is enormously important. From this one equation we will compute all distances and volumes, light travel times, etc. in cosmology.

With the metric in hand, the next step is to derive the equations governing the evolution of the scale factor, $a(t)$. There are three key equations, although only two are unique (the third is derived by combining the first two).

We will begin by deriving the **Friedmann equation**, which is a statement of energy conservation. The proper derivation involves general relativity (GR) but we can derive nearly the entire equation from simple arguments.

Consider a patch of the universe of radius R , undergoing Hubble expansion with velocity v , and with energy per unit mass E . We assume that as this patch of the universe expands it conserves energy. In general we know that the total energy (per unit mass) is the sum of the kinetic and potential energy, so $E = K + U = v^2/2 - GM/R = C$. Let's re-write this equation in comoving coordinates by setting $R = ar$ and $v = \dot{a}r$ where r is the comoving coordinate. Upon substitution we have:

$$E = \frac{(\dot{a}r)^2}{2} - \frac{GM}{ar} = C \quad (2.10)$$

setting $M = 4/3\pi R^3\rho$ we have:

$$E = \frac{(\dot{a}r)^2}{2} - \frac{4}{3}\pi(ar)^2\rho G = C \quad (2.11)$$

Rearranging:

$$\left(\frac{\dot{a}}{a}\right)^2 - \frac{8\pi G\rho}{3} = \frac{C}{a^2} \quad (2.12)$$

The correct form of the Friedmann equation, which includes the effects of general relativity, is:

$$\boxed{\left(\frac{\dot{a}}{a}\right)^2 - \frac{8\pi G\rho}{3} = \frac{-k c^2}{R_c^2 a^2}} \quad (2.13)$$

where R_c is the radius of curvature of the universe, and $k = -1, 0, +1$ for negatively curved, flat, and positively curved universes. In the above, ρ includes both matter and energy densities. Eq. 2.13 is the first of our key equations for describing the homogeneous universe.

We can define the **critical density**, ρ_c , as the density required to produce a flat universe. Setting $k = 0$ on the right hand side, we find

$$\boxed{\rho_c = \frac{3H^2}{8\pi G}} \quad (2.14)$$

where we have used the fact that $H = \dot{a}/a$. Notice that $H = H(t)$ because $a = a(t)$. At the present epoch, $\rho_c^0 \approx 2.78 \times 10^{11} h^2 M_\odot \text{Mpc}^{-3} \approx 5 \times 10^9 \text{eV m}^{-3}$.

“Little h”: You may have noticed the h unit in the numerical expression for ρ_c . For a long time the Hubble constant was not known to better than a factor of two ($50 - 100 \text{ km s}^{-1} \text{Mpc}^{-1}$); for this reason the Hubble constant is often written as $H_0 = 100 h \text{ km s}^{-1} \text{Mpc}^{-1}$. That way, the dependence of any number or expression that depends on H_0 can be easily identified, and the user can easily scale a result to their preferred value of H_0 . Nowadays, with the value of H_0 known to a few percent, this convention is less necessary and in fact often leads to confusion.

The second important equation is known as the **fluid equation** (for cosmology) and it represents mass-energy conservation; its derivation can be sketched in the following way. The expansion of the universe as a whole is adiabatic (i.e., no net heat flow into or out of the universe), which implies $dQ = d\mathcal{E} + PdV = 0$, where $\mathcal{E}$ is the internal energy. We then have:

$$\frac{d\mathcal{E}}{dt} = -P \frac{dV}{dt} \quad (2.15)$$

with $V = \frac{4}{3}\pi a^3$ and $\mathcal{E} = \rho c^2 \frac{4}{3}\pi a^3$. So:

$$\frac{d(\rho c^2 a^3)}{dt} = -P \frac{da^3}{dt} \quad (2.16)$$

apply chain rule:

$$\dot{\rho} c^2 a^3 + 3\rho c^2 a^2 \dot{a} + 3Pa^2 \dot{a} = 0 \quad (2.17)$$

and simplify:

$$a\dot{\rho} c^2 + 3\dot{a}(\rho c^2 + P) = 0 \quad (2.18)$$

$$\boxed{\dot{\rho} + 3\frac{\dot{a}}{a}\left(\rho + \frac{P}{c^2}\right) = 0} \quad (2.19)$$

The third equation, the **acceleration equation**, can be derived by combining the Friedmann and fluid equations. We begin by taking a time derivative of the Friedmann equation:

$$\dot{a}^2 - \frac{8\pi G}{3}\rho a^2 = C \quad (2.20)$$

$$2\dot{a}\ddot{a} - \frac{8\pi G}{3}\dot{\rho}a^2 - \frac{8\pi G}{3}\rho 2a\dot{a} = 0 \quad (2.21)$$

simplifying:

$$\ddot{a} = \frac{8\pi G}{3}\rho a + \frac{4\pi G}{3}\frac{\dot{\rho}a^2}{\dot{a}} \quad (2.22)$$

now we can use the fluid equation to remove the $\dot{a}$ factor from above, leaving us with:

$$\boxed{\frac{\ddot{a}}{a} = -\frac{4\pi G}{3}\left(\rho + \frac{3P}{c^2}\right)} \quad (2.23)$$

Essentially all of homogeneous cosmology is based on these three equations.

Some comments are in order. First, the Friedmann equation tells us that the geometry of the universe (open, flat, or closed) is directly connected with the density of the universe. So a measurement of the latter enables a measurement of the former. Much effort has been expended over the past several decades to measure the density of the universe in order to infer its geometry. Notice also that $\ddot{a}$ is positive (i.e., acceleration) if $P < -\rho c^2/3$. Negative pressure is obviously very weird.

In order to understand the fate of the universe (acceleration, deceleration and collapse, etc) we need to know about the **equation of state** (EOS) of the matter in the universe, i.e., we need to know $P = P(\rho)$. In cosmology it is often useful to express the EOS as $P = w\rho c^2$ where w is a parameter that depends on the nature of the fluid. For example, for photons $w = 1/3$, while for pressureless non-relativistic matter $w = 0$. The latter can be derived by noting that $P = \rho kT/m$ (ideal gas), so:

$$w = \frac{P}{\rho c^2} = \frac{\rho kT}{m\rho c^2} \propto \frac{\langle v^2 \rangle}{c^2} \approx 0 \quad (\text{non-relativistic matter}) \quad (2.24)$$

Cosmic acceleration occurs when $w < -1/3$. Einstein's famous **cosmological constant**, Λ (see below), is equivalent to $w = -1$. Inserting this into the fluid equation (Eq. 2.19) implies that $\dot{\rho}_\Lambda = 0$. Think about how weird that is.

Let's return to the fluid equation and set $P = w\rho c^2$. This leads to:

$$\dot{\rho} + 3\frac{\dot{a}}{a}(1+w)\rho = 0 \quad (2.25)$$

$$\frac{d\rho}{\rho} = -3(1+w) \frac{da}{a} \quad (2.26)$$

$$\rho(a) = \rho_0 a^{-3(1+w)} \quad (2.27)$$

examples:

$$w = 0 \rightarrow \rho \propto a^{-3} \quad (2.28)$$

$$w = 1/3 \rightarrow \rho \propto a^{-4} \quad (2.29)$$

$$w = -1 \rightarrow \rho = \text{const.} \quad (2.30)$$

Recall that $\rho_c = \frac{3H^2}{8\pi G}$. Now let's define the density of an arbitrary component i in units of the critical density: $\Omega_i \equiv \rho_i/\rho_c$. This allows us to write the Friedmann equation in a useful form:

$$\frac{H^2}{H_0^2} = \Omega_r^0 a^{-4} + \Omega_m^0 a^{-3} + \Omega_k^0 a^{-2} + \Omega_\Lambda^0 \quad (2.31)$$

where variables with a “0” subscript/superscript are present-day values, and r and m refer to radiation and matter. Finally, $\Omega_k = 1 - \Omega_r - \Omega_m - \Omega_\Lambda$; for a flat universe $\Omega_k = 0$.

Cosmological constant (Λ): Einstein preferred a static universe ($\dot{a} = \ddot{a} = 0$), but GR naturally produces an evolving universe due to the radiation and matter content. What to do? Einstein decided to add an *ad hoc* term to the Friedmann equation:

$$H^2 = \frac{8\pi G\rho}{3} - \frac{kc^2}{R_c^2 a^2} + \frac{\Lambda}{3} \quad (2.32)$$

Equivalently, one can define $\rho_\Lambda = \frac{\Lambda}{8\pi G}$ and incorporate Λ into the density term. The acceleration equation then becomes:

$$\frac{\ddot{a}}{a} = -\frac{4\pi G}{3} \left(\rho + \frac{3P}{c^2} \right) + \frac{\Lambda}{3} \quad (2.33)$$

and so even in a universe with matter, for a certain value of Λ one can have $\ddot{a} = 0$. As we saw above, if Λ is constant, then $\dot{\rho}_\Lambda = 0$. This is obviously unusual. One way to interpret this is as a constant vacuum energy density. We might appeal to quantum mechanics to estimate the vacuum energy density via $\epsilon_\Lambda \sim \frac{E_p}{l_p^3} \sim \frac{m_p c^2}{l_p^3}$ where m_p and l_p are the Planck mass and length, respectively: $m_p \equiv (\hbar c/G)^{1/2} = 10^{-8}$ kg and $l_p \equiv (G\hbar/c^3)^{1/2} = 10^{-35}$ m. The problem is that this estimate is 124 orders of magnitude larger than the present-day critical density. The origin of the cosmological constant is one of the most important unsolved problems in science.

We now turn to the task of computing cosmological quantities, including distances, volumes, times, etc. This task would all be straightforward if the universe was static. However, the time-dependent expansion makes this problem more challenging, both computationally and also conceptually. For example, due to the finite light travel time and the expansion of space, when considering the separation between objects we must also specify a reference epoch. Hogg (2000) provides a concise pedagogical introduction to this topic and is recommended for further reading. See Fig. 2.1 for relations between various distances that will be computed below.

We return to the Robertson-Walker metric, focusing on variation along the line-of-sight, so that $d\Omega = 0$, and focusing further on light ($ds^2 = 0$), so that the metric reduces to the simple expression:

$$dr = \frac{c dt}{a} \quad (2.34)$$

where dr is a differential length element in *comoving* units.

We can turn this into a more useful expression by considering several useful identities: $da = -(1+z)^{-2}dz = -a^2dz$, and $dt = \frac{da}{aH}$; combining these results in: $dt/a = -dz/H$. This allows us to write:

$$\frac{dr}{dz} = \frac{c}{H(z)} \quad (2.35)$$

We can now define the **comoving distance** to redshift z along the line of sight as:

$$D_C(z) = \int dr = \int_0^z \frac{c}{H(z')} dz' \quad (2.36)$$

The **physical (proper) distance** is the distance measured between two points at a fixed $a(t)$:

$$D_p = \frac{D_C}{1+z} \quad (2.37)$$

To gain some intuition for these different distances, see Fig. 2.1. Consider galaxy B indicated by the second red dashed line. Light from this galaxy that was emitted when the universe was 5 Gyr old ($z \approx 1.3$) would be reaching us now. At that epoch the physical distance between us and galaxy B was ≈ 1.75 Gpc (labeled r_e in the figure). The corresponding comoving distance is $1.75(1+1.3) = 4$ Gpc. This is also the physical distance between us and galaxy B at the present epoch (labeled r_o in the figure).

The next quantity of interest is the **angular diameter distance**. Imagine at a time t_e an

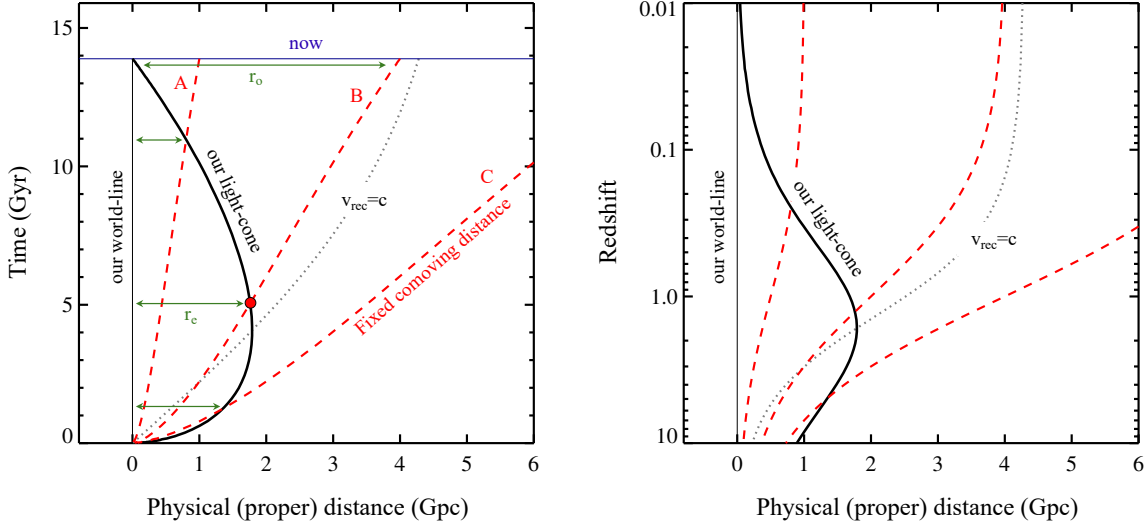


Figure 2.1: Spacetime diagram for the concordance cosmology. Panels show physical (proper) distance vs. time (left) and redshift (right). Our world-line is the thin vertical solid line. The thick solid line is our past light-cone, and the dotted line marks the physical distance at which the recession velocity is equal to the speed of light. Red dashed lines show the physical distance of objects at comoving distances of 1, 4, and 8 Gpc (lines A, B, and C). Consider object B. It is at a comoving (and hence present physical distance) of 4 Gpc. We would observe light from object B when the universe was 5 Gyr old (at $z \approx 1.3$) as indicated by the red dot in the left panel. At that epoch the physical distance to object B was 1.75 Gpc (denoted r_e in the left panel). The physical distance to an object at the epoch the light was emitted is equivalent to the angular diameter distance. Notice that this distance grows with increasing redshift until $z \approx 1.5$; at higher redshifts the angular diameter distance decreases. The origin of this strange behavior can be seen by comparing the solid black and gray dotted lines. Starting at the present epoch, the past light cone grows in the usual way you're familiar with from special relativity. At earlier epochs the expansion velocity at the physical distance of our light cone approaches the speed of light; this causes the rate of change of the light cone physical distance to decline. Eventually the recession velocity along the past light cone exceeds the speed of light and the light cone contracts toward earlier epochs.

object has an angular extent $d\theta$ and a physical length dl . Then, by the RW metric:

$$dl = a(t_e) S_k(r) d\theta \quad (2.38)$$

$$= (1+z)^{-1} S_k(r) d\theta \quad (2.39)$$

We can now define the angular diameter distance, D_A , as the ratio of an object's physical

transverse size to its angular size (in radians):

$$D_A \equiv \frac{dl}{d\theta} = \frac{S_k(r)}{1+z} \quad (2.40)$$

Notice that the angular diameter distance is sensitive to the curvature of the universe. In the case of a flat universe we have the simple expression: $S_k(r) = r = D_C$. We can therefore understand the angular diameter distance as simply the (transverse) proper distance when the light was emitted from the source, i.e., $D_A = (1+z)^{-1}D_C = D_p(z)$. This distance has the unusual property of being non-monotonic with redshift — the relation between D_A and z turns over at $z \gtrsim 1$ (see Fig. 2.1). This means that objects of a fixed size will appear *larger* on the sky at $z \gtrsim 1$ compared to $z \lesssim 1$. We can see where this unusual behavior comes from by differentiating D_A with respect to redshift:

$$\frac{dD_A}{dz} = \frac{1}{1+z} \frac{dD_C}{dz} - \frac{D_C}{(1+z)^2} = \frac{c}{(1+z)H} - \frac{D_A}{1+z} = \frac{c - D_A H}{(1+z)H} \quad (2.41)$$

Notice that the derivative is zero when $c = D_A H$, where $D_A H$ is the expansion velocity at the angular diameter distance. In essence, the angular diameter distance has a maximum where the expansion velocity equals the speed of light. See the caption to Fig. 2.1 for details.

In a flat, static universe, the observed flux at Earth, F , is related to the luminosity L and distance r by:

$$F = \frac{L}{4\pi r^2} \quad (2.42)$$

In an expanding universe with curvature a more general relation is required. We define the **luminosity distance**, D_L :

$$D_L(z) \equiv \sqrt{\frac{L}{4\pi F}} \quad (2.43)$$

for a source with intrinsic luminosity L at cosmological redshift z .

If we consider bolometric quantities (i.e., flux and luminosity integrated over all wavelengths), then the observed flux is the observed power (energy per unit time) per area, while the intrinsic luminosity is the emitted power. At a given physical distance r between the source and Earth, the emitted photons are spread over an area $4\pi S_k(r)^2$ for a universe with arbitrary curvature. We can also relate the emitted and observed wavelengths, unit times, and scale factors as:

$$\frac{\lambda_{\text{em}}}{\lambda_{\text{obs}}} = \frac{\delta t_{\text{em}}}{\delta t_{\text{obs}}} = \frac{a_{\text{em}}}{a_{\text{obs}}} \quad (2.44)$$

hence the observed power can be related to the emitted power via:

$$P_{\text{obs}} = \frac{h\nu_{\text{obs}}}{\delta t_{\text{obs}}} = \frac{h\nu_{\text{em}}}{\delta t_{\text{em}}} \frac{a_{\text{em}}^2}{a_{\text{obs}}^2} = P_{\text{em}} \frac{a_{\text{em}}^2}{a_{\text{obs}}^2} \quad (2.45)$$

and hence the observed flux is:

$$F = \frac{P_{\text{obs}}}{4\pi S_k(r)^2} = \frac{L}{4\pi S_k(r)^2} \frac{a_{\text{em}}^2}{a_{\text{obs}}^2} = \frac{L}{4\pi S_k(r)^2} \frac{1}{(1+z)^2} \quad (2.46)$$

where in the last equality we set $a_{\text{obs}} = 1$ assuming observations are taking place at the present epoch. Relating the expression above to our definition of D_L in Eq. 2.43, we see that:

$$D_L = (1+z) S_k(r) \quad (2.47)$$

Finally, if we compare to Eq. 2.40, we have:

$$D_L = (1+z)^2 D_A \quad (2.48)$$

where the above expression holds for a universe with any degree of curvature. In a flat universe recall that $S_k(r) = D_C$ and so in such a universe $D_L = (1+z) D_C$. A reminder that expressions above are for bolometric quantities; consideration of monochromatic flux entails additional factors (known as the k-correction).

The surface brightness of an object is the flux per unit solid angle: $I \equiv F/\Omega$. In a static universe I is *constant*. However, in an expanding universe this is not the case. To see why, let's define

$$\Omega = \frac{A}{4\pi D_A^2} \quad (2.49)$$

for some object of physical area A . We then have:

$$I = \frac{L}{4\pi D_L^2} \frac{4\pi D_A^2}{A} \quad (2.50)$$

$$= \frac{L}{A} (1+z)^{-4} \quad (2.51)$$

$$= I_e (1+z)^{-4} \quad (2.52)$$

where I_e is the intrinsic surface brightness of the source. The key point is that surface brightness is a *very* strong function of redshift: surface brightness observations at high redshift are therefore very challenging. In 1930 Tolman proposed measuring the surface brightness of galaxies as a function of redshift as a test of whether the universe was static or expanding.

Question: Why might the Tolman test not be a strong test of the nature of the universe?

Recall that the Hubble parameter can be written as $H = \dot{a}/a$, and after some manipulation we can write:

$$dt = \frac{1}{H(z)} \frac{da}{a} = -\frac{1}{H(z)} \frac{dz}{(1+z)} \quad (2.53)$$

so that the **lookback time** to redshift z can be computed via:

$$t_{\text{lb}}(z) = \int_0^z \frac{dz'}{H(z')(1+z')} \quad (2.54)$$

and similarly, the **age of the universe** can be computed via:

$$t_{\text{univ}}(z) = \int_z^\infty \frac{dz'}{H(z')(1+z')} \quad (2.55)$$

or

$$t_{\text{univ}}(z) = t_H \int_z^\infty \frac{dz'}{(1+z') \sqrt{\Omega_r^0(1+z')^4 + \Omega_m^0(1+z')^3 + \Omega_k^0(1+z')^2 + \Omega_\Lambda^0}} \quad (2.56)$$

where $t_H \equiv H_0^{-1}$ is the Hubble time.

In the next chapter we will put these various expressions to good use.

Chapter 3

Cosmological Tests

In the previous chapter we assembled the key equations relevant for describing the state and evolution of a homogeneous universe. In this chapter we will put those equations to use in describing specific limiting cases, as well as how we go about measuring cosmological parameters. We will also introduce the cosmic distance ladder and the concordance cosmological model.

The key equations from the previous chapter are:

$$ds^2 = -c^2 dt^2 + a(t)^2 [dr^2 + S_k^2(r) d\Omega^2] \quad (\text{RW})$$

$$H^2 - \frac{8\pi G\rho}{3} = \frac{-k c^2}{R_c^2 a^2} \quad (\text{I})$$

$$\dot{\rho} + 3 H \left(\rho + \frac{P}{c^2} \right) = 0 \quad (\text{II})$$

$$\frac{\ddot{a}}{a} = -\frac{4\pi G}{3} \left(\rho + \frac{3P}{c^2} \right) \quad (\text{III})$$

$$P = w\rho c^2 \quad (\text{EOS})$$

Recall that $H(a) = \frac{\dot{a}}{a}$ and combining the EOS and Equation II leads directly to $\rho(a) = \rho_0 a^{-3(1+w)}$. We will now work through some classic example limiting cases:

- **Einstein-de Sitter space:** Here we assume a flat, matter-dominated universe: $\Omega_{\text{tot}} = \Omega_m = 1$ and $w = 0$. We then have

$$\rho = \rho_0 a^{-3} \quad (3.1)$$

this should be obvious - the density of normal matter declines in proportion to the volume. Combining Equation I with the above leads to:

$$(\dot{a}/a)^2 = 8\pi G \rho_0 a^{-3}/3 \quad (3.2)$$

$$\dot{a}/a \propto a^{-3/2} \quad (3.3)$$

which leads to:

$$a = \left(\frac{t}{t_0} \right)^{2/3} \quad (3.4)$$

this is the first time that we have been able to write the scale factor as an explicit (and simple) function of time. Notice that with $a(t)$ now specified, we can determine everything else of interest in this toy universe - the evolution of the Hubble parameter, the $t(z)$ relation, distance vs. redshift, etc.

- **Radiation dominated universe:** Here we assume $\Omega_{\text{tot}} = \Omega_r = 1$ and $w = 1/3$. In this case we have

$$\rho_r = \rho_0 a^{-4} \quad (3.5)$$

this is less intuitive, but the idea is that we get three powers of a from the volume term and an additional power of a because the wavelengths of the photons are being stretched along with the expansion of space, reducing their energy. Recall further from basic thermodynamics that the energy density of a blackbody obeys $\rho_r \propto T^4$ where T is the temperature. Therefore, for radiation we also have: $T \propto a^{-1}$.

We can work through Equation I in this case and we arrive at:

$$a = \left(\frac{t}{t_0} \right)^{1/2} \quad (3.6)$$

- **Λ -dominated universe:** Here we assume $\Omega_{\text{tot}} = \Omega_\Lambda = 1$ and $w = -1$. In this case $\rho_\Lambda = \text{constant}$, which as we remarked upon last chapter is very weird. Appealing again to Equation I, we have:

$$(\dot{a}/a)^2 = 8\pi G \rho_\Lambda/3 \quad (3.7)$$

Because the RHS is constant, and the LHS is equal to H^2 , we have: $H_0^2 = 8\pi G\rho_\Lambda/3 = \text{constant}$. Then:

$$\dot{a}/a = H_0 \quad (3.8)$$

$$\int d\ln a = \int H_0 dt \quad (3.9)$$

$$\boxed{a = e^{H_0(t-t_0)}} \quad (3.10)$$

The key conclusion is that a Λ -dominated universe (or any universe with a constant energy density) undergoes **exponential expansion**. This same principle underlies the exponential expansion predicted to occur during inflation.

- **Milne model (empty universe):** Here we assume an empty universe in which $\rho = 0$. In this case Friedmann’s equation (Equation I) has the form: $\dot{a}^2 = -kc^2/R_c^2$. For $k = 0$ we have $\dot{a} = 0$, i.e., a flat and static universe. For $k = -1$ we have $a \propto t$, sometimes called a “coasting universe”. Finally, the $k = +1$ solution is not allowed – one cannot have an empty, positively curved universe. According to GR, positive curvature requires the existence of an energy density.

The real universe contains matter (both dark and baryonic), radiation, and (we believe) Λ . Moreover, because $\rho_r \propto a^{-4}$, $\rho_m \propto a^{-3}$, and $\rho_\Lambda = \text{const}$, we can generally expect any particular epoch to be well-described by being dominated by a single component, so we will often talk about matter-dominated, radiation-dominated, or Λ -dominated epochs (see Fig. 3.1). The limiting cases discussed above are therefore in fact very useful, as they describe the behavior of our universe during particular epochs. The epoch at which $\rho_r = \rho_m$ is known as the epoch of **matter-radiation equality**.

Much of modern cosmology is focused on measuring **cosmological parameters**. Several of the most important are the densities of matter (which includes dark matter), radiation, and dark energy (Λ). These are most frequently parameterized by Ω_m , Ω_r , and Ω_Λ . These quantities are almost always assumed to be the present-day values, with the understanding that scaling to other redshifts is trivial.

In the previous chapter we noted that one useful way to express the Friedmann equation is as follows:

$$\frac{H(z)^2}{H_0^2} = \Omega_r (1+z)^4 + \Omega_m (1+z)^3 + \Omega_k (1+z)^2 + \Omega_\Lambda \quad (3.11)$$

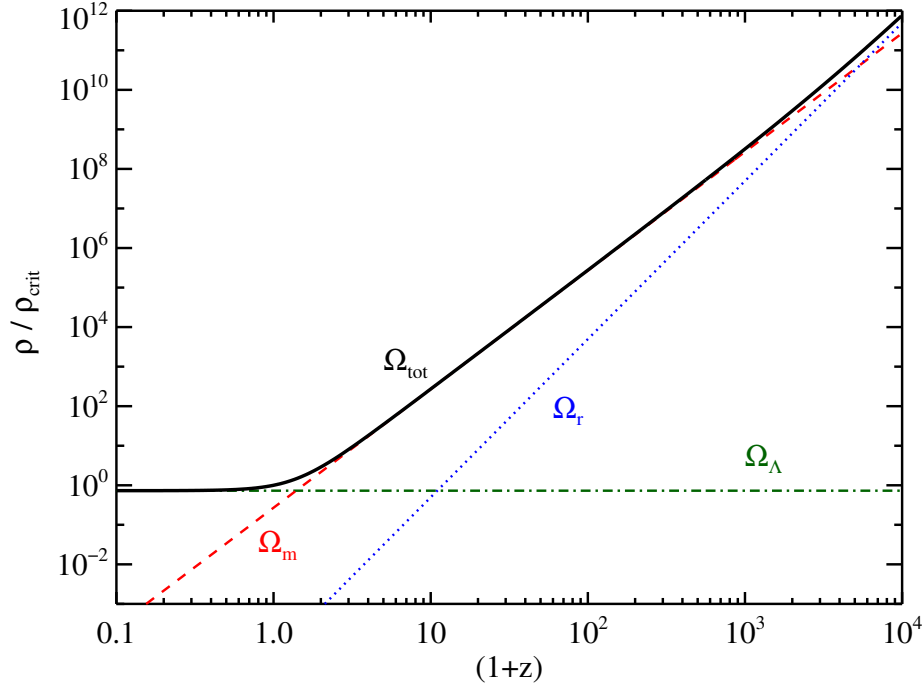


Figure 3.1: Density of the universe in units of the present-day critical density, including the total density and the three individual components. A concordance, flat, Λ CDM model is assumed. Notice that at most epochs the density evolution is well-approximated by that of the dominant component at that epoch. The future evolution is at $(1+z) < 1$.

In the previous chapter we also expressed the comoving distance as an integral of the Hubble parameter:

$$D_C = D_H \int_0^z \frac{H_0}{H(z')} dz' \quad (3.12)$$

where $D_H \equiv c/H_0$ is the Hubble length.

To take one simple example, consider a flat two-component universe containing matter and Λ . In this universe the distance-redshift relation would be:

$$D_C = D_H \int_0^z [\Omega_m (1+z')^3 + \Omega_\Lambda]^{-1/2} dz' \quad (3.13)$$

If we could measure distances to objects as a function of redshift then we could infer the values of Ω_m and Ω_Λ . Because the equation above is an *integral* over these quantities, we require distances measured at multiple redshifts to disentangle these (and other) parameters. Note also that the Hubble constant, H_0 sets the overall length scale and must either be known by independent means or simultaneously constrained. Of course, one can turn the problem

around: if the cosmological parameters are well-determined by other means, then we can use the above equation to estimate distances given redshifts.

In addition to the energy content of the universe, the equation of state of dark energy is another key cosmological parameter. This is often simply referred to as “ w ” and how close this parameter is to -1.0 is obviously of fundamental importance. Even tiny deviations from -1.0 would imply that dark energy is *not* a constant and could have deep implications for our understanding of this mysterious component of our universe. The Dark Energy Task Force (Albrecht et al. 2006) was a major community effort to compare various observational strategies for understanding the nature of dark energy. The task force road map lays out all of the major surveys now ongoing and planned, including BOSS, DES and PFS, LSST, DESI, SKA, and WFIRST. That’s impressive given that the report was written nearly 20 years ago.

Much of the work in cosmological parameter determination boils down to the task of determining distances (comoving distances, angular diameter distances, luminosity distances) to objects as a function of redshift. This is a difficult problem.

How do we measure distances to astronomical objects? The story is fairly complicated, and we often talk about the **distance ladder** as we move from more accurate to less accurate means of measuring distances. Within the solar system we determine distances directly via radar. Beyond the solar system we use parallax, which is another direct measurement of distance that requires no knowledge of the nature of the source. Thanks to the *Gaia* satellite we now have parallax distances for some 2 billion sources to distances of several kpc.

Megamasers are another geometric distance technique (in addition to parallax and radar) that plays an important role in the distance ladder. These are galaxies with an accretion disk around a supermassive black hole in which maser emission occurs within the accretion disk. The majority of megamasers are OH or H₂O megamasers; this is important because they are very bright at radio frequencies, which allows for superb angular resolution via interferometry. One can see how a distance can be derived from a megamaser in the following way. Consider a blob of masing material in orbit around a BH. The physical distance between the BH and the blob is r , and the distance from Earth to the blob is then $d = r/\Delta\theta$ where $\Delta\theta$ is the angular separation. Because these are gaseous blobs, we can assume circular orbits, so that the acceleration is $a = v^2/r$. Then the distance to the blob is $d = \frac{v^2}{a\Delta\theta}$. All of the quantities on the right hand side

are, at least in principle, directly observable. The above assumes that we are viewing the disk edge-on, but with enough information we can also solve for the inclination of (and even the presence of warps in) the disk. Because the timescales are short and the velocities are large, we can directly measure both the velocity and acceleration of the blob, allowing a precision measurement of the distance. The disk sizes are small (~ 0.1 pc), necessitating subarcsecond resolution for galaxies at Mpc distances. By far the best-studied megamaser is NGC 4258 at a distance of 7.2 Mpc. This one object is a very important anchor in the extragalactic distance scale.

Beyond these two techniques, distances to essentially all other objects in the universe are indirect and rely on properties of the source. The classic example is the **standard candle**, which is an object whose intrinsic luminosity is known and whose observed flux can then be used to determine a distance. The trick of course is to identify objects whose luminosity is known and stable with very high precision and accuracy. An early example of a standard candle is the Cepheid variable. Working at the Harvard College Observatory around the turn of the last century, Henrietta Leavitt studied Cepheids in the Large Magellanic Cloud and discovered a strong correlation between their periods and magnitudes. Periods can be measured for other Cepheids at unknown distances, and the period-luminosity relation can be used to infer the luminosity, producing a standard candle. Cepheids were used by Hubble to conclude that the universe is expanding, and Cepheids are still the primary means for measuring the Hubble constant. The main limitations are possible effects of metallicity on the P-L relation, and the effect of dust on the observed fluxes.

Cepheids have luminosities of $\sim 10^{3.5} - 10^{4.5} L_{\odot}$, which means it becomes increasingly difficult to detect them further away. For reference, a typical Cepheid has an absolute V -band magnitude of about -6.0 . It is generally difficult to detect anything, even with *HST* fainter than about 29th magnitude. This sets a distance limit for Cepheids of about 100 Mpc, or $z \approx 0.025$.

The most luminous objects in the universe (in the optical) are the supernovae (SNe). They are not generally standard candles, but the Type Ia SNe, which are believed to be the thermonuclear explosion of a white dwarf, are known to be *standardizable*. The light curves of the Ia's have regular behavior, such that there is a strong correlation between the peak brightness and the width of the light curve, known as the Phillips Relation (see Fig. 3.2). This results in Ia SNe being effectively standard candles (the physical origin is still an active area of research, but is likely due to the uniformity in the synthesized Ni-56 mass, which powers the light curve). The shape of the light curve does not itself produce a standard candle; for that one needs to tie the Cepheids and the SNe together. This is done for the

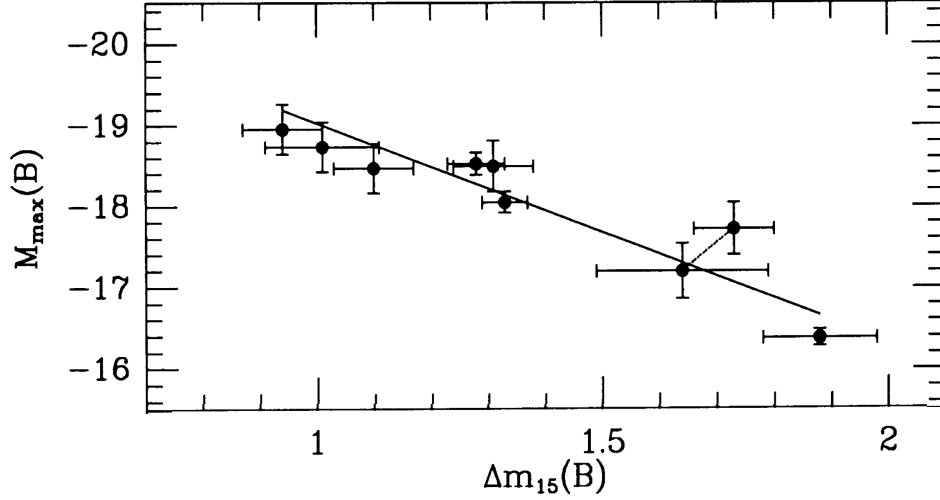


Figure 3.2: Relation between SN Type Ia peak brightness and the change in brightness over the first 15 days following the peak; from Phillips (1993). This is the first evidence that Type Ia SNe are “standardizable”. The x-axis can be measured for any SN, enabling determination of the intrinsic SN brightness.

small number of galaxies (19) for which both Type Ia SNe and Cepheids have been identified.

Type Ia SNe have been the workhorse in cosmology for the past 30 years. In the late 90s measurements to $z \sim 0.5$ indicated cosmic acceleration (see Fig. 3.3). That result led to a Nobel Prize. The challenges with SNe as standard candles are legion, including the fact that the physical origin is not well-known, so it is entirely possible, if not likely that the luminosity could depend on unknown variables related to the host population. Dust is also a source of concern, as are flux zero points and k-corrections. Everything in astronomy is hard at the 1% level.

Standard candles provide a measurement of the luminosity distance, D_L . **Standard rulers** provide a measurement of the angular diameter distance, D_A . Attempts were made for years to use galaxies as both standard candles and standard rulers, but the reality is that galaxies evolve too strongly with redshift to be used for precision cosmology. The two most important standard rulers are the cosmic microwave background (CMB) and the baryon acoustic oscillation (BAO). Both are due to the same basic physics, namely the distance a sound wave travels in the radiation fluid before recombination. In the case of the CMB, we have a ruler at $z \sim 1000$ while in the case of the BAO we have many rulers in the redshift range of $0.3 \lesssim z \lesssim 2$. The BAO scale is ≈ 100 Mpc, which means that very large volumes are needed in order to provide a strong statistical measurement.

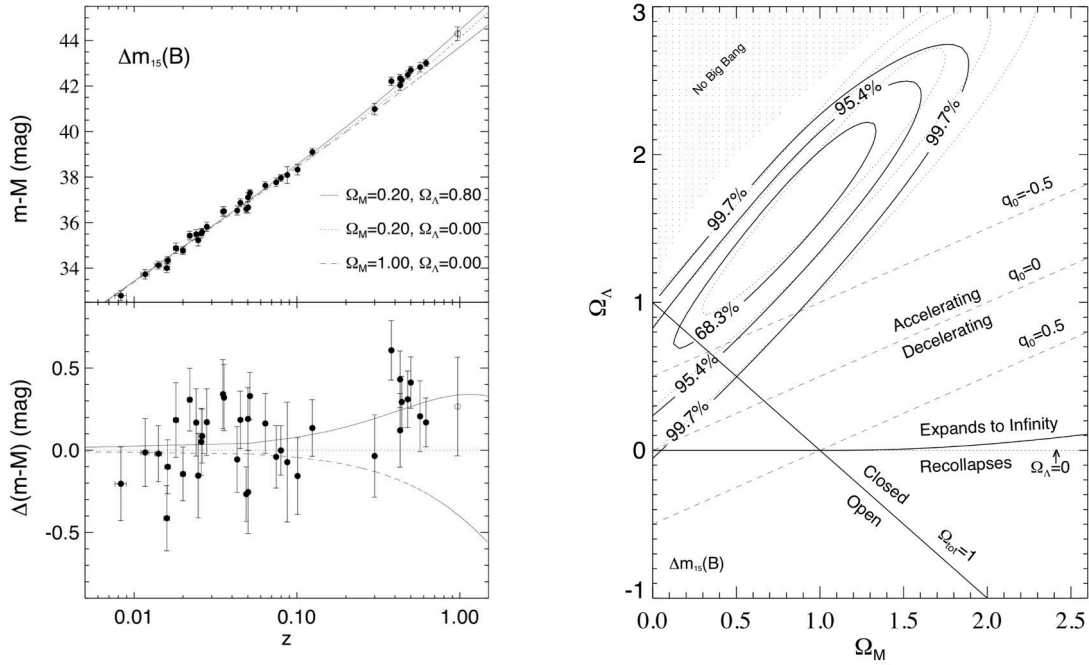


Figure 3.3: Evidence for cosmic acceleration. The left panels show the distance moduli of type Ia SNe vs. redshift compared to three cosmological models: Einstein-de Sitter ($\Omega_{\text{tot}} = \Omega_m = 1$), an open universe ($\Omega_{\text{tot}} = \Omega_m = 0.2$) and a flat, Λ dominated universe ($\Omega_{\text{tot}} = \Omega_m + \Omega_\Lambda = 0.2 + 0.8$). The right panel shows a fit to the SNe data in the space of $\Omega_\Lambda - \Omega_m$, showing 1, 2, 3 σ contours. From Riess et al. (1998).

There are also **standard shapes**, i.e., something expected to be spherical (either individually or statistically). This includes the BAO and voids. The angular direction provides an integral constraint on $H(z)$ as we've seen before. The shape along the line of sight provides a local constraint on $H(z)$ because $dz/dr = H(z)/c$. Enforcing a particular shape, e.g. a sphere, provides a constraint on the ratio of $D_A/H(z)$.

The final classic cosmological test comes from **age constraints**. The basic idea here is to find the oldest objects in the universe and use them to rule out models that predict younger ages. This is all quite model-dependent, as nearly all ages depend on models. The two main clocks used in this context are main sequence turnoff ages and white dwarf cooling ages. In the mid-90s there was a puzzle in that stellar physics was providing ages of the oldest stars of ≈ 15 Gyr, older than the emerging consensus cosmological models (10 – 13 Gyr). It turned out that the ages were wrong, due to a combination of inaccurate distances and various subtle issues in the stellar models (opacities and abundance effects). Ages are no longer a competitive constraint on cosmology (although some groups still use the age evolution of

the oldest galaxies in the universe as a constraint on $H(z)$), but it is worth emphasizing how remarkable it is that two completely disjoint approaches both yield ages of the universe consistent at the $\approx 10\%$ level.

Of all the techniques, the BAO is probably the most promising and reliable at the accuracy required for current and future constraints. In later weeks we will discuss constraints on cosmology emerging from the growth of perturbations (e.g., weak lensing and cluster counts).

Cosmology underwent a dramatic revolution around the turn of the century. Prior to this revolution, there was widespread disagreement about the basic nature of the universe - whether the universe was flat, open, or closed, whether or not dark energy existed, whether or not inflation existed, etc. etc. In the late 1990s and early 2000s two research programs came to fruition: 1) the SNe programs clearly favored dark energy and $\Omega \approx 1$; 2) the WMAP CMB experiment very strongly favored a flat universe containing a mix of dark matter and dark energy. This convergence to a standard, fiducial model for the universe is frequently referred to as **concordance cosmology**. The basic tenets of the concordance cosmology are that the universe is flat $\Omega = 1.0$ and dominated by dark energy and dark matter: $\Omega_\Lambda \approx 0.70$ and $\Omega_m \approx 0.30$. The contribution of baryons to the total density is $\Omega_b \approx 0.05$, implying a baryon fraction of $\Omega_b/\Omega_m \approx 0.17$ - that is, 17% of all matter is in the form of baryons. The present-day radiation density is $\Omega_r \approx 5 \times 10^{-5}$. There are other aspects of the concordance cosmology related to recombination, inflation, and the nature of primordial density fluctuations that we will discuss in later chapters.

The convergence to a fiducial cosmological model has been enormously helpful to the field of galaxy formation. In the 1990s through early 2000s it was common to see results in papers quoted assuming wildly different cosmologies, or multiple simulations run with very different cosmologies, because the uncertainties in cosmology were an important component of the error budget. Now, the residual uncertainties in the cosmological model are a negligible source of uncertainty in our understanding of galaxies.

Of course, all of this is subject to the very important caveat that the nature of dark matter and especially dark energy are completely mysterious. We must therefore leave open the possibility that our entire cosmological model, however successful, is in need of major revision.

Chapter 4

The Cosmic Microwave Background

In this chapter we will discuss the cosmic microwave background (CMB), including the history and basic physics. Key concepts include recombination and decoupling (and understanding the difference between those concepts), the baryon-to-photon ratio, the Saha equation, and the evolution of the ionization fraction with redshift.

The history of the CMB is fascinating. In a series of papers in 1946–1948, Gamow, Bethe, and Alpher predicted the existence of the CMB based on big bang nucleosynthesis calculations. In order to explain the light element abundances, they appealed to primordial temperatures of at least $T \sim 10^{10}$ K (comparable to $m_e c^2$), which implied a present-day radiation temperature of ≈ 5 K. This result was far ahead of its time. Twenty years later, Dicke, Peebles, and Wilkinson¹ (1965) from Princeton revisited the connection between an early hot universe and subsequent expansion, resulting in a present-day radiation temperature of ~ 10 K (as with many cases in cosmology, Soviet scientists independently came to similar conclusions at nearly identical times; in this case it was Doroshkevich & Novikov in 1964). At exactly the same time, and completely independently, Penzias and Wilson, working at Bell Labs in New Jersey, noticed background noise in a new large radio antenna, which ended up being the serendipitous discovery of the CMB. For this discovery, the two received the Nobel Prize in 1978.

Following the discovery of the CMB and subsequent theoretical efforts, there was a wide range of observational efforts. The most consequential was COBE, launched as a satellite in 1989 (and resulted in Nobel Prizes for the PIs in 2006). COBE confirmed that the CMB is basically a perfect blackbody with $T \approx 3$ K and temperature fluctuations at the level of

¹Wilkinson is the “W” in the WMAP acronym.

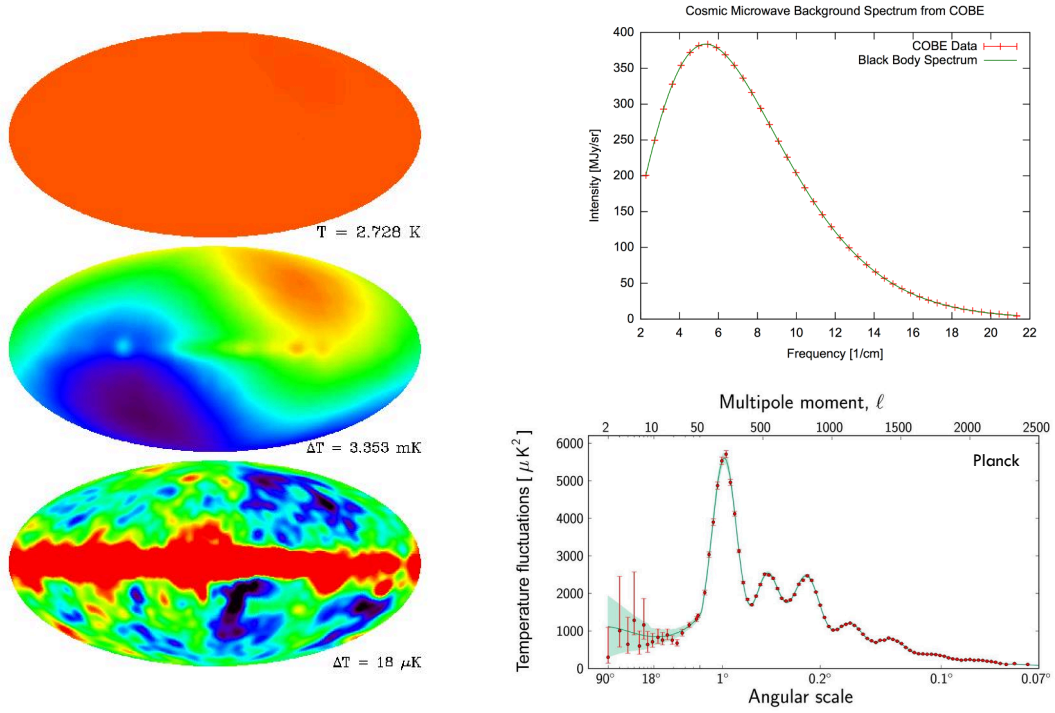


Figure 4.1: Left panels: All-sky temperature maps from the COBE satellite. The top panel shows the temperature with a fairly wide temperature range for the color map. The middle panel shows a much narrower temperature range to highlight the dipole imprinted from the Earth’s motion around the Sun. The bottom panel shows the temperature fluctuations after removing the dipole signature. The bright stripe in the middle is the Milky Way galaxy. Top right panel: intensity vs. frequency from COBE data compared to a blackbody with $T = 2.73\text{K}$. Bottom right panel: temperature power spectrum from the *Planck* satellite. The acoustic scale is clearly visible at $\approx 1^\circ$.

$|\Delta T/T| \sim 10^{-5}$. There is also a dipole in the CMB due to the Earth’s motion at the level of $|\Delta T/T| \sim 10^{-3}$. This dipole is almost always removed from the “final” CMB map that is usually presented. See Fig. 4.1.

More recently, the *WMAP* and *Planck* satellites (launched by NASA in 2003 and ESA in 2013, respectively) provided much more precise measurements of the CMB, ushering in the era of “precision cosmology” (see Fig. 4.2).

There are several important basic features of the CMB:

- the CMB is isotropic to very high precision (after removing the expected dipole due to the Earth’s motion).
- the CMB is a nearly perfect blackbody.

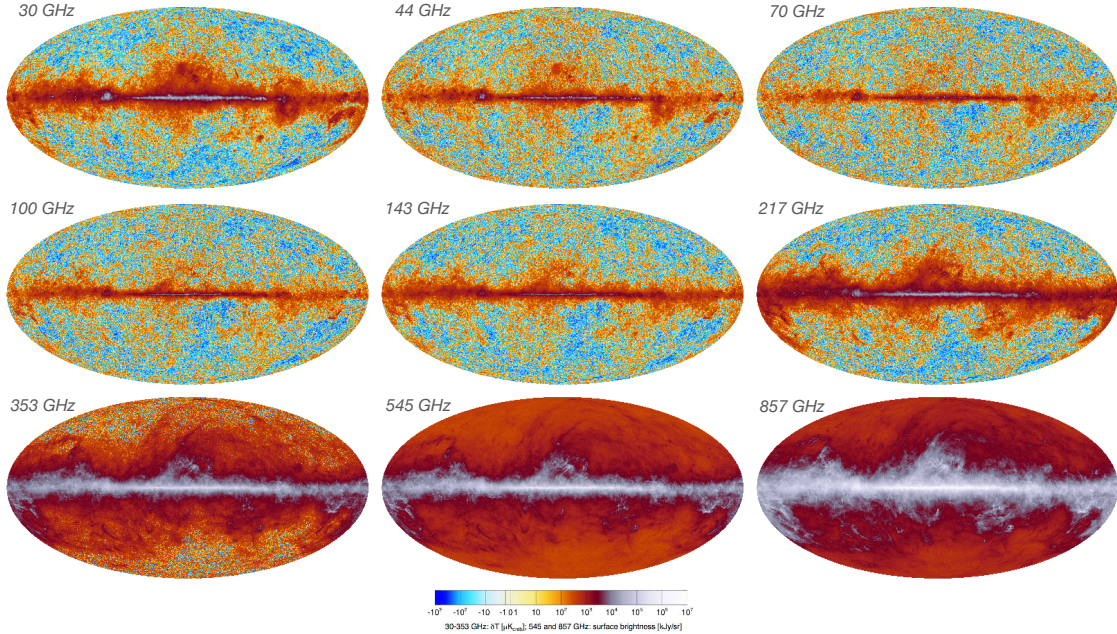


Figure 4.2: All-sky intensity maps from the *Planck* satellite at different frequencies. Notice the very different structure at different frequencies. High frequencies are affected by synchrotron and free-free radiation, while low frequencies are affected by thermal dust emission from the Galaxy (figure from www.cosmos.esa.int/web/planck/picture-gallery).

- the CMB is not resolved into sources even at arcsecond resolution (this sounds like an obvious statement given our theoretical prejudice, but it is important to separate what the data alone tells us from what data plus theory imply).
- the fact that the CMB is a blackbody means that the CMB was emitted from an optically-thick source. We observe radio and optical sources at $z > 5$, implying that the optical depth to $z \sim 5$ must be small. This in turn implies that the CMB was emitted at $z > 5$.
- there is a lot of energy and an enormous number of photons in the CMB.

An important concept in the context of the CMB is the **baryon-to-photon ratio**, η , which we can compute as follows. The radiation density in the CMB today is: $\rho_r = \alpha T^4 = 0.26 \text{ MeV m}^{-3}$ today (α is the radiation constant). The mean energy per photon is quite small: $h\langle\nu\rangle \approx 6.3 \times 10^{-4} \text{ eV}$, implying that the number density of photons in the CMB is: $n_\gamma = \rho_r/h\langle\nu\rangle \approx 4 \times 10^8 \text{ m}^{-3}$. Let's compare this to the baryon density. We know from a variety of observations that $\Omega_b = 0.05$, so $\rho_b = 0.05\rho_c = 0.05 \cdot 5200 \text{ MeV m}^{-3} = 260 \text{ MeV m}^{-3}$. This implies $\rho_b/\rho_r \approx 1000$, i.e., the energy density of the baryons is much higher than that

of radiation. But: the rest energy per baryon is quite high, about 940 MeV (i.e., the rest energy of the proton). So the number density of baryons is only $n_b = \rho_b/940 \text{ MeV} \approx 0.28 \text{ m}^{-3}$. Hence the baryon to photon ratio is:

$$\boxed{\eta \equiv n_b/n_\gamma = 7 \times 10^{-10}} \quad (4.1)$$

or approximately 1.4×10^9 photons per baryon. This large value of η has important consequences for our understanding of the evolution of the universe.

We will now discuss the concepts of **recombination** and **decoupling**.

Define the **fractional ionization**, X , as:

$$X \equiv \frac{n_p}{n_p + n_H} = \frac{n_p}{n_b} = \frac{n_e}{n_b} \quad (4.2)$$

where n_p , n_H , and n_e are the number densities of protons, neutral hydrogen atoms, and electrons. The last equality arises from the assumption of charge neutrality in the universe (i.e., the number of protons and electrons are the same). Recall that the process of ionization of H involves:

$$H + \gamma \leftrightarrow p + e^- \quad ; \quad Q = 13.6 \text{ eV} \quad (4.3)$$

furthermore, matter and radiation are coupled via Compton scattering:

$$\gamma + e^- \leftrightarrow \gamma + e^- \quad (4.4)$$

which, in the limit of low electron velocities ($v_e^2/c^2 \ll 1$) is called Thomson scattering, where the Thomson cross section is $\sigma_e = 7 \times 10^{-25} \text{ cm}^2$.

A simple (and incorrect) estimate of decoupling:

The mean free path (mfp) for a particle undergoing collisions (or scatterings) is $\lambda = (n\sigma)^{-1}$ where n is the number density of colliders and σ is the cross section for interaction. The scattering *rate* is then $\Gamma = v/\lambda$ for particles moving with velocity v . In our case we are interested in the mfp of photons colliding against a sea of electrons, so the scattering rate is:

$$\Gamma = \frac{c}{\lambda} = c n_e \sigma_e = \frac{c n_b^0 \sigma_e}{a^3} \quad (4.5)$$

where the last equality arises because we have assumed that the universe is completely ionized, so that $n_e = n_p = n_b$ and the density of baryons at any epoch is $n_b(a) = n_b^0 a^{-3}$, where a is the scale factor.

We are now going to compare the scattering rate to the Hubble parameter, H . Both of these quantities have units of $[\text{time}]^{-1}$ and can both be thought of as a rate (or an inverse

timescale). The epoch when $H > \Gamma$ is called **freeze-out** or **decoupling**, and it is when the scattering timescale is longer than the age of the universe. As we saw in Chapter 1, when a process has a timescale longer than the age of the universe, that process is not important. Before decoupling, the photons and baryons will be in thermodynamic equilibrium, and the photons will have high optical depth. After decoupling the baryons and photons can have different temperatures, and the optical depth plummets - the universe becomes transparent.

If H remained ionized, then decoupling would occur at $z = 36$. At this redshift the CMB has a temperature of $T = 100$ K. Of course, at such a low temperature hydrogen is no longer ionized. What actually happens is that the energy of the photons drops as a^{-1} until hydrogen recombines². As it is the free electrons that are scattering with photons, the scattering rate will rapidly drop when recombination occurs, resulting in much earlier decoupling than we have computed here. *We therefore need to first compute the recombination of electrons and protons before we can accurately estimate the decoupling time.* We spent time on this incorrect version of decoupling because it allowed us to lay out the basic physical processes.

The simplest estimate of recombination:

We can attempt a naive estimate of when recombination occurs as follows. Assume that recombination occurs when $kT_{\text{CMB}} \approx 13.6$ eV. This would be a reasonable first guess, as that is the energy required to ionize hydrogen. According to the Planck function, the mean energy per photon is $2.7kT$. Therefore, we can estimate the temperature at recombination as $T_{\text{rec}} = \frac{13.6 \text{ eV}}{2.7k} \approx 60,000$ K.

A less-simple estimate of recombination:

While a good first guess, this estimate is badly incorrect because of the enormous number of photons per baryon ($\eta^{-1} \sim 10^9$), which means that there are plenty of high energy photons in the Wien tail of the blackbody to keep hydrogen ionized. We can obtain a somewhat better estimate by letting the Boltzmann suppression factor (which governs the Wien tail of the Planck function) be $e^{-Q/2.7kT} \sim \eta$, which implies, for $Q = 13.6$ eV and $\eta = 7 \times 10^{-10}$, $T \approx 2770$ K. This is still quite crude, but much closer to the correct answer.

A better approach to recombination:

A better approach is to use the Saha equation in an expanding universe to compute the evolution of the ionization fraction, X , with time. Recall that the Saha equation relates the number densities of particles involved in an ionization-recombination process, for example

²The “re” in recombination is a misnomer – this is the first time in the evolution of the universe in which free protons and electrons combine to form hydrogen.

hydrogen, protons and electrons:

$$\frac{n_H}{n_p n_e} = \left(\frac{m_e kT}{2\pi\hbar^2} \right)^{-3/2} e^{Q/kT} \quad (4.6)$$

where $Q = 13.6$ eV. Our goal is to transform this equation into one that is a function of X , T , and η . We can use our earlier expression for X to set $n_H = \frac{1-X}{X} n_p$ and we know that $n_p = n_e$, so:

$$\frac{1-X}{X} = n_p \left(\frac{m_e kT}{2\pi\hbar^2} \right)^{-3/2} e^{Q/kT} \quad (4.7)$$

next, set $\eta = n_b/n_\gamma = \frac{n_p}{X n_\gamma}$. From the properties of a blackbody we have:

$$n_\gamma = \frac{2.4}{\pi^2} \left(\frac{kT}{\hbar c} \right)^3 \quad (4.8)$$

combining the previous few expressions yields: $n_p = 0.24X\eta(kT/\hbar c)^3$. Putting everything together:

$$\boxed{\frac{1-X}{X^2} = 3.84\eta \left(\frac{kT}{m_e c^2} \right)^{3/2} e^{Q/kT}} \quad (4.9)$$

From this equation we can compute $X(t)$ because we know for the radiation $T \propto a(t)$ and the cosmological model provides $a(t)$. We can define **recombination** as when $X = 0.5$. For our concordance cosmology this implies $T_{\text{rec}} = 3700$ K and $t = 240,000$ yr ($z = 1370$). Recombination is not instantaneous, though it is fairly rapid: $X = 0.9$ occurs at $z = 1475$ and $X = 0.1$ occurs at $z = 1255$, for a time interval of $\Delta t = 70,000$ yr, or, in terms of the age of the universe at that epoch: $\Delta t/t = 70,000/240,000 = 30\%$. See Fig. 4.3.

Decoupling done right:

Now that we have the evolution of the ionization fraction with time, we can return to our calculation of decoupling:

$$\Gamma = n_e(z)\sigma_e c = X(z)(1+z)^3 n_b^0 \sigma_e c \quad (4.10)$$

assuming $\Omega_b^0 = 0.05$ we have:

$$\Gamma = 5.5 \times 10^{-21} \text{ s}^{-1} X(z) (1+z)^3 \quad (4.11)$$

during the recombination epoch the universe is matter-dominated ($\rho_m > \rho_r$) so the evolution of the Hubble parameter is reasonably well-approximated as:

$$\frac{H^2}{H_0^2} = \Omega_m^0 (1+z)^3 \quad (4.12)$$

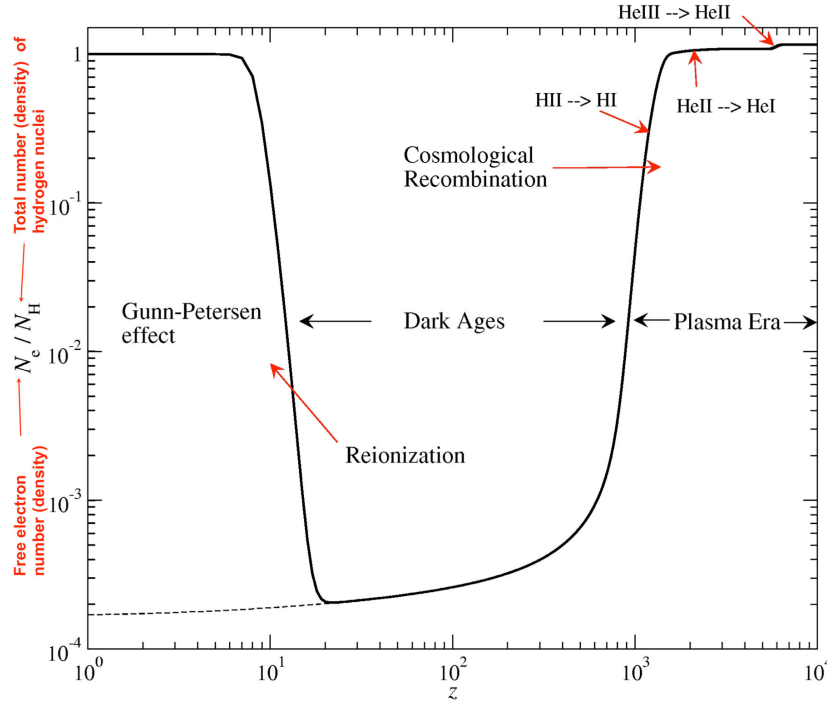


Figure 4.3: Free electron number density vs. redshift, normalized to the total number density of hydrogen atoms. The transitions associated with recombination and reionization are labeled. From Sunyaev & Chluba (2009).

as before, decoupling occurs when $\Gamma = H$, which leads to:

$$1 + z_{\text{dec}} = \frac{37}{X(z_{\text{dec}})^{2/3}} \quad (4.13)$$

which, after solving a messy equation, leads us to: $z_{\text{dec}} = 1120$ and $X(z_{\text{dec}}) = 0.006$.

Notice that while recombination and decoupling occur at roughly the same epoch, they refer to two distinct processes, and decoupling occurs only after the majority (99%) of the protons and electrons have combined.

The probability of scattering is $dP = \Gamma dt$. We can compute the expected number of scatterings between some initial time t and today (t_0) as:

$$\tau(t) = \int_t^{t_0} \Gamma(t) dt \quad (4.14)$$

where τ is the optical depth. The epoch t when $\tau = 1$ is known as the **last scattering surface**, t_{ls} . This is analogous to the surface of a star. In practice $t_{\text{ls}} \approx t_{\text{dec}}$. The CMB is the last scattering surface.

Even in the absence of ionizing photons, recombination takes time and proceeds according to the following equation:

$$\frac{dn_e}{dt} = -\alpha n_e n_p = -\frac{n_p}{t_{\text{rec}}} \quad (4.15)$$

where α is the recombination coefficient and is equal to $\alpha = 3 \times 10^{-13} (T/10^4 \text{ K})^{-0.8} \text{ cm}^3 \text{ s}^{-1}$ for the conditions of interest, and t_{rec} is the timescale for recombination. As we saw before, this process will freeze-out when $t_{\text{rec}} = H^{-1}$. The math is left to the reader; the result is a final residual electron fraction of $X_{\text{final}} \sim 10^{-4}$.

Let's summarize what we've learned. At early times ($z \gg 1000$) the universe is fully ionized and Thomson scattering results in strong coupling between the photon and baryon fluids. The expansion of the universe produces a declining baryon (electron) density and a declining radiation temperature. Eventually the temperature drops to the point that there are an insufficient number of high energy photons to keep hydrogen ionized. As protons and electrons begin to (re)combine, the electron fraction drops, which in turn results in a decrease of the scattering rate between the baryons and the photons (see Fig. 4.3). This gives rise to a decoupling between these two fluids, and allows the photons to freely stream throughout the universe. This is the origin of the CMB, and is sometimes known as the surface of last scattering.

We have sketched out the key physical processes governing the interaction between baryons and photons and the subsequent emergence of the CMB. Detailed calculations are far more complicated, although the basic physical processes and ingredients are well-understood.

Chapter 5

Cosmic Inflation

In the first four chapters we have presented the standard model of cosmology and noted its success in explaining a wide variety of observations. In this chapter we will discuss several of its shortcomings, and a proposed solution: inflation.

The standard model of cosmology suffers from three problems: 1) the flatness problem; 2) the horizon problem; 3) the monopole problem. The first two are relatively straightforward to understand and are genuine problems in need of a solution. The third, the monopole problem, emerges from theoretical high energy particle physics.

The **flatness problem** can be understood from manipulation of the Friedmann equation, which can be written as:

$$1 - \Omega(t) = -\frac{kc^2}{R_c^2 a(t)^2 H(t)^2} \quad (5.1)$$

where $\Omega(t)$ is the total energy density of the universe (excluding the Ω_k term introduced in Chapter 2). At the present epoch we have: $1 - \Omega_0 = -kc^2/R_c^2/H_0^2$. Combining these two expressions leads to:

$$1 - \Omega(t) = (1 - \Omega_0) \frac{H_0^2}{H(t)^2 a(t)^2} \quad (5.2)$$

If we consider only radiation and matter, then the Hubble parameter changes with scale factor according to:

$$\frac{H(t)^2}{H_0^2} = \Omega_{r,0}a^{-4} + \Omega_{m,0}a^{-3} \quad (5.3)$$

so:

$$1 - \Omega(t) = (1 - \Omega_0) \frac{a^2}{\Omega_{r,0} + \Omega_{m,0}a} \quad (5.4)$$

The significance of the above equation is that the universe *becomes less flat* as we go back in time. Adopting the concordance cosmological parameters of $\Omega_{r,0} = 5 \times 10^{-5}$ and $\Omega_{m,0} = 0.3$ we find that if the universe is flat today to 1% (i.e., $|1 - \Omega_0| < 0.01$), then it must have been flat to a level of 2×10^{-5} at matter-radiation equality, and at the time that the elements were forming (the epoch of Big Bang Nucleosynthesis, at $a \approx 10^{-8}$) the universe must have been flat at the level of one part in 10^{14} . This is the flatness problem: the universe must have been *extremely* flat in the early universe for it to be as flat as we observe today. Nothing in the standard model of cosmology tells us why the universe must be so flat. This can be thought of as a fine-tuning problem: absent some natural explanation, we have to just assert that the universe was perfectly flat at $t = 0$, in spite of the fact that the universe could well have chosen any value of curvature.

The **horizon problem** is a more serious issue, and can be understood as follows. A photon travels a physical distance $cdt = dx = a(t)dr$, so that the particle horizon is:

$$R_H(t) = a(t) \int_0^t \frac{cdt'}{a(t')} \quad (5.5)$$

Recall that for matter-dominated universes $a \propto t^{2/3}$ while for radiation-dominated universes $a \propto t^{1/2}$. By coincidence, this means that $R_H \propto t$ in both cases.

We can estimate the ratio of comoving horizon volumes today and at some earlier epoch (assuming matter-dominated expansion) via:

$$\frac{R_H^3(t_0)/a_0^3}{R_H^3(t)/a(t)^3} \sim (1+z)^{3/2} \quad (5.6)$$

such that at decoupling there were approximately 10,000 causally disconnected volumes. Said another way, the horizon distance at decoupling was only 0.4 Mpc, corresponding to an angular scale of $\approx 2^\circ$.

The horizon problem can be summarized as follows: at the time of decoupling photons could only have traveled about 0.4 Mpc, or $\approx 2^\circ$, and yet the entire sky is uniform in temperature

to a level of 10^{-5} . Because the particle horizon shrinks faster than the scale of the universe itself ($R_H \propto t$ while $a \propto t^{1/2}$), these causally-disconnected regions were *never* in causal contact and so there was no way for these regions to thermally equilibrate. And yet that is precisely what we observe. This is in some ways a deeper problem than the flatness problem, as one requires a very improbable degree of chance coordination among many disconnected regions of the universe.

We will not spend much time on the monopole problem, as it emerges as a theoretical puzzle from ideas in high energy particle physics. The basic concept is that various theories for the unification of the fundamental forces predict that the universe goes through one or more phase transitions as the energy scale of the universe declines, and these transitions are generally expected to lead to *topological defects*¹, including magnetic monopoles. Models predict a density of monopoles that exceeds the matter density of the universe by many orders of magnitude, which is clearly not observed. A period of early inflation reduces the density of these defects by many orders of magnitude, to a level that they would not easily be observed today.

It turns out that a period of exponential expansion of the universe at very early times is sufficient to explain or solve the three problems stated above. This period is known as **inflation**.

Exponential expansion is a natural outcome of the Friedmann equation when the equation of state is $w = -1$. In this case we have $P = -\rho c^2$ and ρ is constant. If we neglect the curvature term for the moment, then we have:

$$H^2 = \frac{8\pi G\rho}{3} \quad (5.7)$$

where H is constant. Because $H = \dot{a}/a$, we have

$$\boxed{a(t) = a_0 e^{Ht}} \quad (5.8)$$

i.e., exponential expansion. Phrased another way:

$$\frac{a(t_f)}{a(t_i)} = e^N \quad , \quad N \equiv H(t_f - t_i) \quad (5.9)$$

¹An example of a topological defect is the structure observed in frozen water, e.g., an ice cube.

where t_i and t_f mark the beginning and end of the inflationary epoch. Typical values of N are ~ 100 , i.e., the universe inflates in size by $\sim 10^{43}$ in the earliest epochs of the universe.

How does inflation solve our problems? Consider first the flatness problem. From above, we have that the flatness can be expressed as:

$$|1 - \Omega(t)| = \frac{c^2}{R_c^2 a(t)^2 H(t)^2} \quad (5.10)$$

During inflation H is constant and $a \propto e^{Ht}$ so that $|1 - \Omega(t)| \propto e^{-2Ht}$. Even if the universe was strongly curved, such that $|1 - \Omega| \sim 1$ before inflation, after $N = 100$ e-foldings the universe would be flat to a level of 10^{-87} .

The horizon problem is solved in a similar manner. The photon horizon during inflation scales as $R_H \propto e^N$, such that after 100 e-foldings the horizon increases by 10^{43} . Various models for unification of the strong and weak forces predict this to occur around an energy scale of 10^{12} TeV, which corresponds to a time in the universe of 10^{-36} s. If inflation began around this time, then the horizon immediately prior to inflation was $\approx 10^{-26}$ cm (the universe was radiation-dominated prior to inflation), and immediately after the horizon was ≈ 1 pc. By the time of decoupling, this scale would have grown to $\sim 10^{43}$ Mpc, obviously large enough to solve the horizon problem.

The magnetic monopole problem is solved in a similar way - inflation essentially dilutes the density of topological defects away to the point that an already rare event would be exponentially more rare in our observable universe after inflation.

Clearly a period of exponential expansion would solve several major problems in the standard cosmological model. But what is inflation exactly, and how does it work (how does it start and why does it end)?

The short answer to these questions is that we do not know. Inflation is really a class of models, not one specific model. This makes inflation difficult to confirm or rule out (which makes the class of models unappealing to some). It is in fact quite remarkable how quickly the inflationary paradigm was adopted in the late 1980s, in light of the scant empirical evidence in support of it at the time. There are at least three basic predictions of inflation: i) a spectral index n of density fluctuations that deviates slightly from 1.0; ii) Gaussian density fluctuations seeded by primordial quantum fluctuations; iii) a tensor to scalar ratio of $r \sim 0.1$. The first prediction has been borne out by observations, with the *Planck* satellite measuring $n = 0.9665 \pm 0.0038$ (we will have more to say about the spectral index in later

chapters). The second prediction is consistent with observations as far as can be measured. The third prediction is a very active area of research (e.g., by John Kovac's group here at Harvard). The tensor to scalar ratio measures the ratio of the power spectra in scalar and tensor modes, where the former are the intuitive density fluctuations and the latter are sourced by primordial gravitational waves.

Here we'll give just a flavor of the physics of inflation. Imagine a scalar field, ϕ , called the **inflaton**, and an associated potential of the field, $V(\phi)$. If the field is homogeneous then we can write the energy density and pressure as:

$$\rho_\phi = \frac{1}{2}\dot{\phi}^2 + V(\phi) \quad (5.11)$$

$$P_\phi = \frac{1}{2}\dot{\phi}^2 - V(\phi) \quad (5.12)$$

where our assumption of homogeneity allowed us to set $\nabla\phi = 0$, and we have adopted $c = \hbar = 1$ (the above equations come from the definition of the energy-momentum tensor in terms of the Lagrangian density of a scalar field).

Notice that if $\dot{\phi}^2 \ll V(\phi)$ then $\rho_\phi = -P_\phi = V(\phi)$, i.e., we have an equation of state $w = -1$. This condition is known as the *slow roll approximation*, and it is a requirement for the inflaton to produce an inflationary epoch.

Inserting Eq. 5.11 into the cosmological fluid equation, $\dot{\rho} + 3H(\rho + P) = 0$, yields:

$$\ddot{\phi} + 3H\dot{\phi} + \frac{dV}{d\phi} = 0 \quad (5.13)$$

This equation is analogous to the equation of motion of a ball moving under the influence of a potential in the presence of friction (Hubble drag). The slow-roll approximation is equivalent to assuming $\ddot{\phi} \ll dV/d\phi$, so that we can neglect the first term on the left hand side of the above equation.

Note that we have two requirements of the potential: that its derivative with respect to ϕ be sufficiently small to produce exponential expansion *and* that its initial amplitude, V_0 , be sufficiently large to dominate the energy density of the universe (see Fig. 5.1). Inflation will begin when the energy density in radiation has dropped below V_0 . Inflation ends when the potential steepens and slow-roll is no longer a good approximation. For the potential shown in Fig. 5.1, the field will oscillate about the minimum, with oscillations damped by Hubble friction. However, we might also expect coupling between the inflaton field and other fields

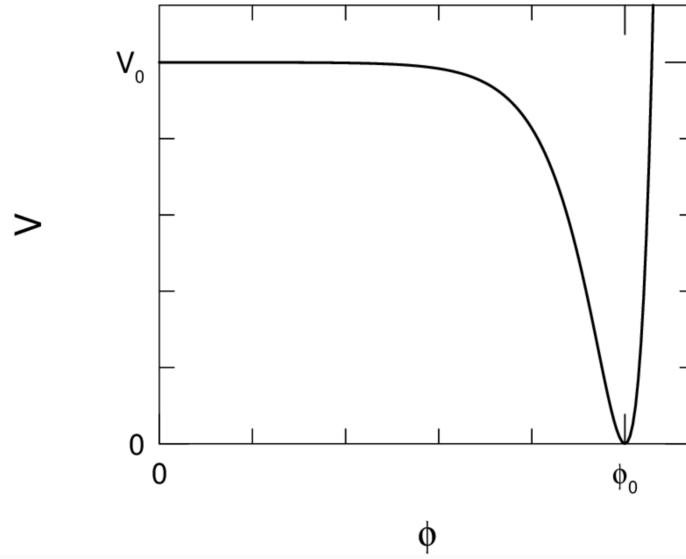


Figure 5.1: A potential that gives rise to an inflationary epoch. The global minimum (“true vacuum”) is at $\phi = \phi_0$. If the scalar field starts at $\phi = 0$ it will slow-roll toward the global minimum, during which time it will give rise to inflation. From Ryden (their Fig 11.3).

in the universe, such as radiation, in which case energy from the inflaton field will likely be transferred to the other fields, resulting in an epoch known as **reheating**. This phase is very uncertain, but it is an essential ingredient in inflationary models because inflation, while succeeding in making the universe flat, will also make it extremely *cold*. If $T \sim 10^{28}$ K immediately prior to inflation, then at the end of inflation $T \sim 10^{-15}$ K. Without some form of reheating, the universe would be a barren wasteland.

The connection, if any, between the inflaton and the cosmological constant, Λ , is unknown. Perhaps $V(\phi_0) = \rho_\Lambda$?

There are many flavors of inflation (early inflation, new inflation, chaotic inflation, stochastic, eternal inflation, etc.). These models are characterized by different forms of the potential, $V(\phi)$. It is worth emphasizing that there is no compelling motivation from particle physics for the existence of the inflaton field nor for the various proposed forms of the potential. In other words, the specific inflationary models are ad hoc.

We’ll conclude with a note that there are alternatives to inflation, the most popular of which is the cyclic model developed by Steinhardt and Turok in 2001. This whole field of early universe cosmology is a very active topic of research. It is also in some ways out at the very boundaries of science - many of the predictions of these cosmological models are not observable, even in principle.

Part III

Gravitational Dynamics

Chapter 6

Potential Theory & Relaxation

In this chapter we will begin our detour into gravitational dynamics by reviewing potential theory and introducing the foundational concept of two-body relaxation.

We can think of the problem of dynamics in terms of the number of bodies involved. The classic two-body problem is analytic and has been studied in enormous detail for centuries. The three-body problem is tremendously difficult and not analytically tractable except in certain restricted forms. The problem of interest for galactic dynamics is the $\sim 10^3\text{--}10^{12}$ -body problem. As you might imagine, the approach to this problem is qualitatively different from the first two, and will be the focus of our efforts in these four chapters.

The primary situation of interest for us is one in which the interactions are *collisionless* (i.e., the only relevant interaction is via gravity). There are many applications of this general setup: cosmological N -body simulations of dark matter, star cluster evolution, galaxy dynamics, planetesimal disks, etc.

Let's start with the basics of potential theory. The force experienced by a body of mass m is $\mathbf{F} = m\mathbf{g}(\mathbf{x})$ where $\mathbf{g}(\mathbf{x})$ is the gravitational field, which is defined as:

$$\mathbf{g}(\mathbf{x}) \equiv G \int d^3\mathbf{x}' \rho(\mathbf{x}') \frac{\mathbf{x}' - \mathbf{x}}{|\mathbf{x}' - \mathbf{x}|^3} \quad (6.1)$$

for an arbitrary mass distribution ρ .

We can define the gravitational potential, Φ , as:

$$\Phi(\mathbf{x}) \equiv -G \int d^3\mathbf{x}' \frac{\rho(\mathbf{x}')}{|\mathbf{x}' - \mathbf{x}|} \quad (6.2)$$

which allows us to write the gravitational field in terms of the gradient of the potential:

$$\mathbf{g}(\mathbf{x}) = -\nabla\Phi \quad (6.3)$$

Mathematical manipulation of the above equations leads eventually (see BT Eq 2.6–2.10) to

Poisson's equation:

$$\nabla^2\Phi = 4\pi G\rho \quad (6.4)$$

This is a differential equation for Φ given ρ , subject to a boundary condition (usually $\Phi \rightarrow 0$ as $x \rightarrow \infty$).

The potential energy, U , of a system is:

$$U = \frac{1}{2} \int d^3\mathbf{x} \rho(\mathbf{x}) \Phi(\mathbf{x}) \quad (6.5)$$

this is the work done against gravity in assembling the mass distribution (recall that in general, the work done against a force in moving a particle from x_1 to x_2 is $W_{12} = -\int_{x_1}^{x_2} d\mathbf{x} \cdot \mathbf{F}$).

Let's consider the simplest case of a point source of mass m . In this case we have:

$$\Phi = -\frac{Gm}{r} \quad ; \quad g = \frac{Gm}{r^2} \quad (6.6)$$

for the potential and the gravitational field at distance r from the point source.

Note that for a general mass distribution, we have $g = \frac{GM(<r)}{r^2}$, but $\Phi \neq \frac{GM(<r)}{r}$, where $M(<r)$, sometimes confusingly written as $M(r)$, is the mass enclosed within a sphere of radius r .

The centripetal acceleration for a circular orbit is v_c^2/r . If we set this equal to g then we can define a circular velocity:

$$v_c = \sqrt{\frac{GM(<r)}{r}} \quad (6.7)$$

We will frequently see this as a way of quantifying the mass distribution.

The energy of each particle is $E = \frac{1}{2}mv^2 + m\Phi(\mathbf{x})$, which is the sum of kinetic and potential energies. Obviously the energy per unit mass is $e = \frac{1}{2}v^2 + \Phi(\mathbf{x})$. The energy of a particle is constant along its orbit if $\Phi(\mathbf{x})$ is constant (in time) along the orbit. Furthermore, if we define $\Phi(x \rightarrow \infty) = 0$, as is common convention, then bound orbits have $E < 0$. The escape speed at a particular location is $v_e = \sqrt{-2\Phi}$.

Escape from a potential well occurs when $e = 0$, i.e., when $v_e^2 = -2\Phi$, where v_e is the **escape velocity**. The average squared escape speed is:

$$\langle v_e^2 \rangle = \frac{\int d^3x \rho v_e^2}{\int d^3x \rho} = \frac{-2 \int d^3x \rho \Phi}{M} = \frac{-4W}{M} \quad (6.8)$$

We will see in Chapter 8 that the virial theorem implies $W = -2K$. Adopting $K = \frac{1}{2}M\langle v^2 \rangle$, we have: $\langle v_e^2 \rangle^{1/2} = 2\langle v^2 \rangle^{1/2}$; the rms escape speed is twice the rms speed. The fraction of particles in a Maxwellian velocity distribution with speeds $> 2\times$ the rms speed is 7.4×10^{-3} .

There are several important (and analytically tractable) potential-density pairs that you should become familiar with. The cases below assume a spherical mass distribution.

- A constant density system, $\rho = C$, has a linearly rising circular velocity: $v_c \propto r$.
- A system with a constant circular velocity has $\rho \propto r^{-2}$ and $\Phi \propto \log r$. Systems with a constant velocity are known as **isothermal spheres**. They have an important role in astronomy as many systems closely approximate an isothermal sphere.
- The Plummer model has a constant density core and a power-law outer component:

$$\Phi = -\frac{GM}{\sqrt{r^2 + b^2}} \quad (6.9)$$

where b is a scale radius. This potential implies a density distribution:

$$\rho(r) = \frac{3M}{4\pi b^3} \left(1 + \frac{r^2}{b^2}\right)^{-5/2} \quad (6.10)$$

which has the limit of $\rho = C$ at small r and $\rho \propto r^{-5}$ at large r . This model is used widely in part because it is analytically tractable.

- The King model builds on the Plummer model but with the addition of a tidal truncation (see Fig. 6.1). The King model is derived from a lowered isothermal distribution function (distribution functions are the subject of the next chapter), which describes a self-gravitating system in quasi-thermal equilibrium, but with a truncation in energy to account for tidal escape. Unfortunately, King models do not have analytic solutions for the density or potential. The tidal radius, r_t , is defined in the King model as the radius where the density drops to zero. The core radius, r_c , (sometimes called the King radius) is the radius where the density has fallen by half. One can then define a concentration as $c \equiv \log_{10}(r_t/r_c)$. The shape of a King model is completely described by c . One often encounters the parameter W_0 as describing a King model. This is the dimensionless central potential: $W_0 \equiv \Phi_0/\sigma^2$. These two options for describing the King model – c and W_0 – are related nearly 1-1; larger W_0 values correspond to more concentrated clusters.
- There are several important double power law models. The Jaffe (1983) model:

$$\rho(r) = \frac{\rho_0}{(r/a)^2 (1 + r/a)^2} \quad (6.11)$$

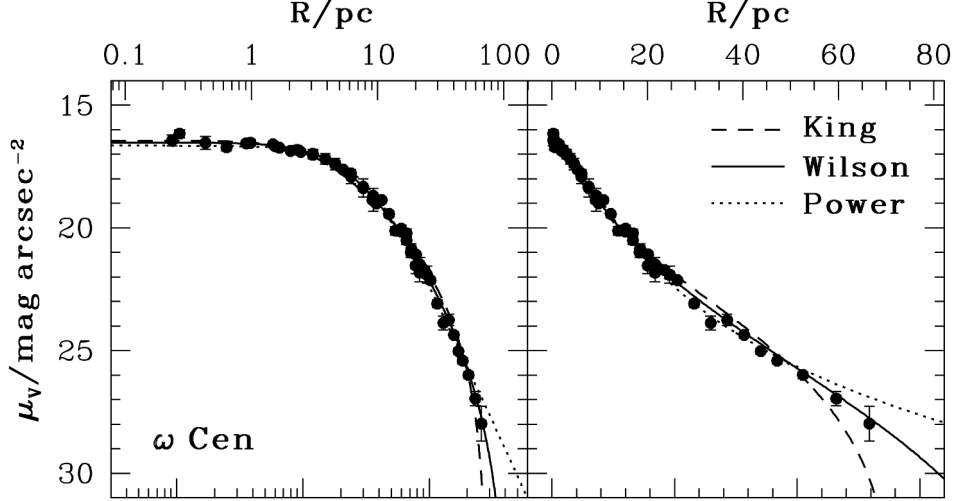


Figure 6.1: Surface brightness profile of omega Cen in log (left) and linear (right) units. The King model does a reasonably good job of describing the light profile for this and other globular clusters. From McLaughlin & van der Marel (2005).

was used to approximate the de Vaucouleurs light profile of elliptical galaxies. The Hernquist (1990) model:

$$\rho(r) = \frac{\rho_0}{(r/a)(1+r/a)^3} \quad (6.12)$$

was proposed as a better approximation to the light profile of elliptical galaxies, and is still widely used for that purpose. The NFW (Navarro, Frenk, and White 1997) model:

$$\rho(r) = \frac{\rho_0}{(r/a)(1+r/a)^2} \quad (6.13)$$

is an excellent approximation to the density profiles of simulated dark matter halos. Notice that the NFW model does not have finite mass.

When the potential is constant, the energy per unit mass of a particle, $e = \frac{1}{2}v^2 + \Phi$, is constant. However, in a time-varying potential, this is no longer the case. We can see this by taking the time derivative of e :

$$\frac{de}{dt} = \frac{\partial e}{\partial v} \frac{dv}{dt} + \frac{\partial e}{\partial \Phi} \frac{d\Phi}{dt} \quad (6.14)$$

$$= -v \nabla \Phi + \frac{d\Phi}{dt} \quad (6.15)$$

$$= -v \nabla \Phi + \frac{\partial \Phi}{\partial t} + v \nabla \Phi \quad (6.16)$$

where the last equation arises from converting the total (Lagrangian) derivative into partials. The result is:

$$\boxed{\frac{de}{dt} = \frac{\partial \Phi}{\partial t}} \quad (6.17)$$

This process of energy exchange caused by a time-varying potential is known as **violent relaxation**.

Notice that the change in energy per unit mass is independent of mass. This is in stark contrast to collisional relaxation (e.g., in a gas, or via two-body effects), where the momentum exchange drives the body toward equipartition of kinetic energies. At the completion of violent relaxation, the particles all will have the same velocity distribution, independent of mass. Violent relaxation is the process by which a collisionless system reaches dynamical equilibrium (virialization).

For a satellite orbiting within a larger host system, an important process by which the satellite loses mass is **tidal stripping**. The basic idea is that the tidal field prunes distant stars from the satellite at a distance from the satellite called the tidal radius, resulting in a loss of mass.

The tidal force exerted by a more massive body M on a less massive body m is:

$$\Delta F = F_{\text{tidal}} = \frac{GMm}{(d-r)^2} - \frac{GMm}{(d+r)^2} \quad (6.18)$$

where d is the distance between the massive body and the center of the less massive body, and r is the distance from the center of the less massive body. In the limit where $r \ll d$ we have:

$$F_{\text{tidal}} = \frac{4GMmr}{d^3} \quad (6.19)$$

the key point is that the tidal force scales as d^{-3} .

The self-gravity force within the smaller body is: $F_{\text{self}} = \frac{Gm^2}{r^2}$ if we set $F_{\text{tidal}} = F_{\text{self}}$ we arrive at the tidal radius:

$$r_t = d (m/4M)^{1/3} \quad (6.20)$$

This is the radius at which material from the smaller body will be removed by tidal forces. This equation gives a good sense of the dependencies and rough order of magnitude, but the details depend sensitively on the orbit of the satellite and the mass distributions of both the satellite and the host. A more complete derivation requires consideration of the Lagrange points in a reference frame co-rotating with the satellite (see BT Ch 8.3.2). The tidal radius has many alternative names including the Jacobi radius, Hill radius, and Roche radius (for example, the star and planet formation community prefer the term Hill radius).

A simple heuristic is to state that the tidal radius occurs where the mean density of the satellite within r_t is equal to four times the mean density of the host inside the orbital radius d : i.e., $\langle \rho_s(r_t) \rangle = 4\langle \rho_h(d) \rangle$.

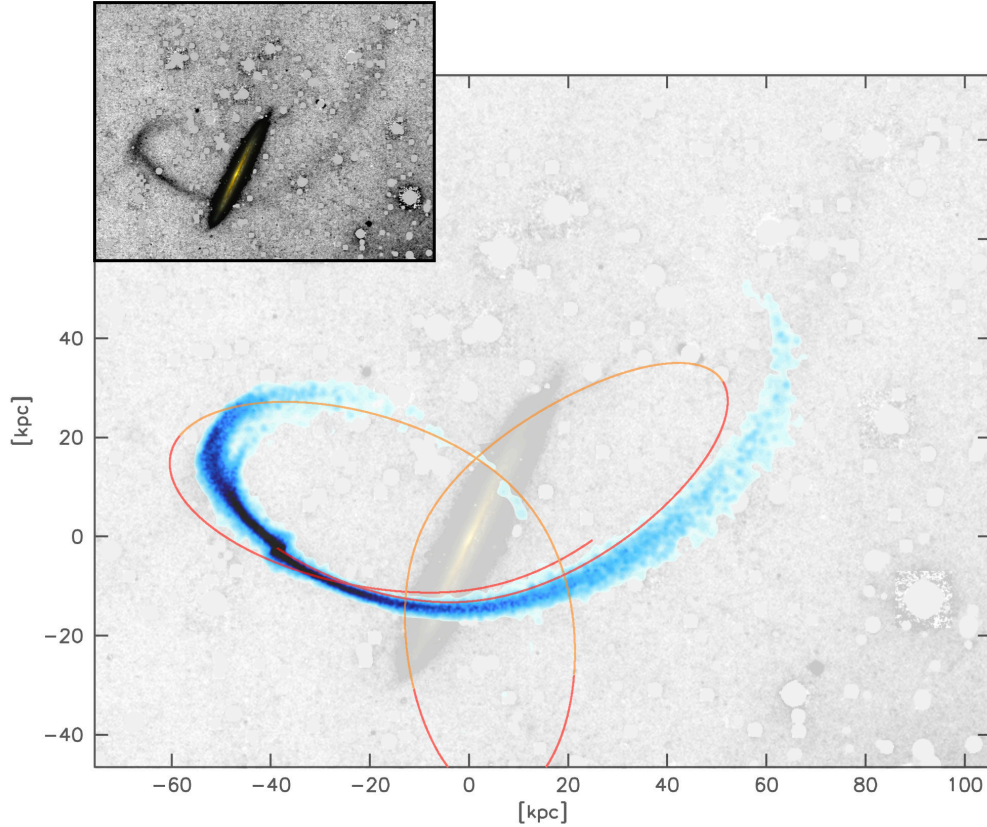


Figure 6.2: Deep imaging of the galaxy NGC 5907 reveals a spectacular stellar stream encircling the galaxy (inset). The observed stream is well-matched to a simple N-body model of an orbiting dwarf galaxy that is being tidally stripped by NGC 5907 (colored density in the large panel). From van Dokkum et al. (2019).

We see evidence for tidal stripping from satellite galaxies around our Galaxy as well as others (see Fig. 6.2). The Sagittarius dwarf is one of the most spectacular examples, as we can see stripped stars wrapping around the entire Galaxy.

We will now consider **two-body relaxation**. Recall that for gas in a box, frequent collisions efficiently redistribute energy, resulting in a Maxwellian velocity distribution ($p \propto e^{-v^2/\sigma^2}$). A stellar system will also relax to a Maxwellian velocity distribution if given sufficient time to do so (although we will see in a later chapter that the thermodynamics of self-gravitating systems is very different compared to the thermodynamics of gases). How long will this take? The basic condition for relaxation is that the energy redistribution process (in this case gravity) must change the stellar velocity by order unity. We will now compute the so-called relaxation time.

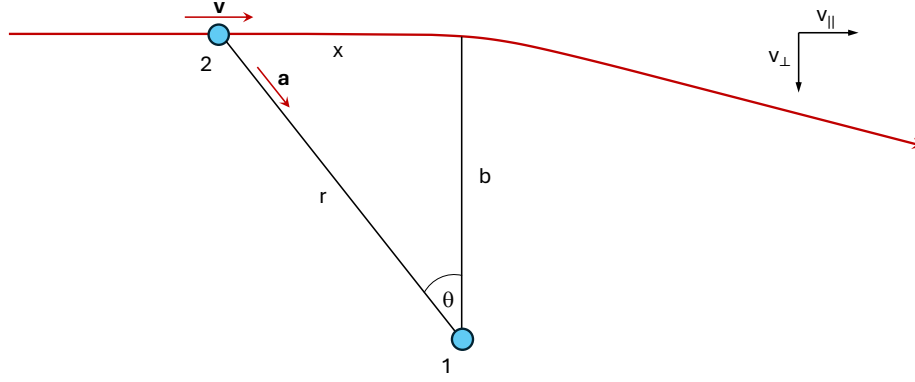


Figure 6.3: Geometry for two-body encounters. Body 2 travels past body 1 at initial velocity v . The impact parameter b is the point of closest approach, r is the distance between the two bodies, and x is the distance between the bodies projected along the direction of motion. After the encounter, body 2 acquires a velocity component perpendicular to the initial direction of motion ($v_{\perp}$).

Consider a two-body interaction with impact parameter b (see Fig. 6.3 for the geometry). The acceleration perpendicular to the velocity vector is:

$$a_{\perp} = \frac{dv_{\perp}}{dt} = \frac{Gm}{b^2 + x^2} \cos\theta = \frac{Gmb}{(b^2 + v^2 t^2)^{3/2}} \quad (6.21)$$

where $a_{\perp} = a \cos\theta$ and $x = vt$. If we now integrate over time throughout the encounter (from $-\infty < t < \infty$) to compute the total change in velocity perpendicular to the velocity vector we have:

$$v_{\perp} = \frac{2Gm}{bv} \quad (6.22)$$

Note that, to first order, there is no change in the velocity parallel to the initial velocity vector by symmetry. The above expression holds in the limit of straight-line trajectories, i.e., when $v_{\perp}/v \ll 1$, which is equivalent to requiring $b \gg 2Gm/v^2$.

Another way to see a rough sketch of the problem is to note that the acceleration at closest approach is $a = Gm/b^2$, and the time to cross the system is $\sim b/v$, so $v_{\perp} \sim a t_{\text{cross}} \sim Gm/bv$.

If we consider many random encounters, we have $\langle v_{\perp} \rangle = 0$, because the encounters are random in direction and orientation. However, and critically, $\sum v_{\perp}^2 \neq 0$ (this is qualitatively analogous to the derivation of a random walk).

In one crossing through the system a star will experience $\delta n = \frac{N}{\pi R^2} 2\pi b db$ interactions with impact parameter in the interval $(b, b + db)$, where N and R are the total number of stars in the system and the system radius.

The total squared change in $v_\perp$ per db for one crossing is:

$$\sum v_\perp^2 = v_\perp^2 \delta n = \left(\frac{2Gm}{bv} \right)^2 \frac{2N}{R^2} b db \quad (6.23)$$

We now integrate over b to obtain the total change in $v_\perp^2$ for one crossing:

$$\Delta v_\perp^2 = 8N \left(\frac{Gm}{vR} \right)^2 \ln \Lambda \quad (6.24)$$

where $\ln \Lambda = \ln \frac{b_{\max}}{b_{\min}}$ is the infamous Coulomb logarithm. The variable Λ quantifies the limits of integration, and hides much of the messiness of the problem. Thankfully, the dependence on Λ only appears in the logarithm, so the details aren't that important. The usual approach is to set $b_{\max} \sim R$, i.e., the largest impact parameter is bounded by the size of the system. The choice for $b_{\min}$ is less well-justified, but is often set to $b_{\min} = 2Gm/v^2$ (this is approximately the impact parameter at which the deflection angle is 90°). The issue at small impact parameter is that the key assumption in the derivation – small changes to the trajectory – breaks down.

Potential theory (or the virial theorem, or simply dimensional arguments) tells us that the velocity will scale as $v^2 \sim GNm/R$. We can then compute the fractional change in the velocity per crossing as:

$$\frac{\Delta v_\perp^2}{v^2} = \frac{8 \ln \Lambda}{N} \quad (6.25)$$

The number of crossings required to change the velocity by of order itself is then $n_{\text{relax}} = \frac{N}{8 \ln \Lambda}$. Therefore, $t_{\text{relax}} = n_{\text{relax}} t_{\text{cross}}$, or:

$$\boxed{t_{\text{relax}} \approx \frac{0.1N}{\ln \Lambda} t_{\text{cross}}} \quad (6.26)$$

Note that

$$\ln \Lambda \approx \ln \frac{R}{b_{\min}} \approx \ln \left(\frac{GNm}{v^2} \frac{v^2}{2Gm} \right) \approx \ln N \quad (6.27)$$

We can write the relaxation time another way by recalling that $t_{\text{cross}} \approx t_{\text{dyn}} = (G\rho)^{-1/2}$. If we also set $\rho = nm \sim Nm/R^3$ (dropping order-unity constants), we have:

$$\boxed{t_{\text{relax}} \approx \frac{0.1v^3}{G^2 \ln \Lambda m^2 n}} \quad (6.28)$$

The relaxation time is the time for a gravitating system to lose memory of its initial conditions. The system will relax to a Maxwellian velocity distribution after t_{relax} . Importantly, if $t_{\text{relax}} \gg t_{\text{univ}}$, then 1) the initial conditions of the system are retained, and 2) the discreteness of the mass distribution is unimportant, i.e., we can approximate the system with a smooth potential. This would allow for enormous simplification of the problem.

Let's consider some example systems:

- **Massive galaxies** have $N \sim 10^{11}$ and $t_{\text{cross}} \sim 10^8$ yr, implying $t_{\text{relax}} \gg t_{\text{univ}}$.
- **Ultra-faint galaxies** have $N \sim 10^4$, $\sigma \sim 5$ km s $^{-1}$, and $R \sim 1$ kpc, so $t_{\text{cross}} \sim 2 \times 10^8$ yr. These systems require care, because the stars are *not* self-gravitating: the velocity dispersion is set by the dark matter, and the stars on their own would produce only $\sigma \approx 0.15$ km s $^{-1}$. The first expression for t_{relax} above therefore cannot be used, as it eliminated v by assuming $v^2 \sim GNm/R$. The second expression keeps v explicit, and evaluating it with the observed σ and the *stellar* number density gives $t_{\text{relax}} \sim 10^{16}$ yr. Two-body relaxation is utterly irrelevant in these systems.
- **Globular clusters** have $N \sim 10^{3-5}$ and $t_{\text{cross}} \sim 10^{5-6}$ yr. They are therefore in a very interesting regime, in which $t_{\text{relax}} \lesssim t_{\text{univ}}$. Some, but not all memory is lost in GCs. In detail, the relaxation time is shorter in the core of the GC, and so we can expect their central regions to be more dynamically evolved than their outer regions. We will discuss the fascinating dynamics of GCs in a later chapter.

Stellar velocities are distributed into a Maxwellian after a relaxation time. The fraction of particles in a Maxwellian velocity distribution with speeds exceeding the escape speed is 7.4×10^{-3} . The process of **evaporation** therefore removes a fraction $f = 7.4 \times 10^{-3}$ of all stars in a system per relaxation time. Evaporation removes stars from a system at a rate:

$$\frac{dN}{dt} = -\frac{7 \times 10^{-3} N}{t_{\text{relax}}} \equiv -\frac{N}{t_{\text{evap}}} \quad (6.29)$$

where $t_{\text{evap}} \approx 140 t_{\text{relax}}$.

Finally, we can compute the timescale for direct stellar **collisions** in a system with a number density of stars n and characteristic velocity v . The collision timescale is approximately $t_{\text{coll}} = (n\sigma v)^{-1}$ where σ is the cross section for collision that will roughly be $\sigma_* = \pi(2r_*)^2$, i.e., the geometric cross section, and r_* is the stellar radius. If we set $v_*^2 = 2Gm/r_*$ as the

escape speed of the star, and approximate $t_{\text{cross}} \sim R/v$ then we have:

$$\frac{t_{\text{coll}}}{t_{\text{cross}}} \approx 0.02N \left(\frac{v_*}{v} \right)^4 \quad (6.30)$$

implying that $t_{\text{coll}} \gg t_{\text{cross}}$ in almost all situations (note that the escape speed for the Sun is $\approx 600 \text{ km s}^{-1}$). Stellar collisions are believed to occur with some frequency in the central cores of some globular clusters.

Chapter 7

Distribution Functions and the Jeans Equations

In the previous chapter we derived the two-body relaxation time and concluded that for most systems of interest (in particular, galaxies) the relaxation time is much longer than the age of the universe. This fact has several consequences, perhaps the most important being that we may now treat each body (star) as an individual orbit within a *smooth* potential. To consider how useful this simplification is, consider the case of globular clusters in which the relaxation time is shorter than the age of the universe. For such systems the graininess of the potential is an important component of the evolution, and so in order to study globular clusters we must model the dynamics of *all* the stars in the system.

When the potential can be assumed smooth, we can take two approaches to describing the system: 1) Consider individual orbits, and build a system by superposition of many orbits (that is consistent with a given potential). This is a computational problem and is known as Schwarzschild modeling. 2) Consider distributions in phase space. In this chapter we will explore the second option.

A warning: this is a very dense chapter, probably best consumed in multiple sittings. Below I mostly follow the notation in BT. An important exception to this is the energy, which BT define such that bound orbits have $E > 0$. As always, use caution when comparing equations from different sources.

Before we tackle phase space, let's remind ourselves of the basic equations of fluid dynamics (Euler's equations). Assume a fluid has density $\rho(\mathbf{x})$ and velocity vector $\mathbf{v}(\mathbf{x})$ at a 3D position

x. Conservation of mass requires:

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{v}) = 0 \quad (7.1)$$

which can be written as:

$$\frac{D\rho}{Dt} = 0 \quad (7.2)$$

where the latter is the total, or convective, or Lagrangian derivative. It is the derivative considered *along* the flow of the fluid (i.e., in Lagrangian space). If we consider Eulerian space (the first equation above), then mass conservation is simply a statement that the change in density at a given point is related to the advection of material toward or away from that point. In Lagrangian space, where we move along with the fluid, the density remains unchanged with time (this is basically the definition of mass conservation).

Conservation of linear momentum in the fluid results in the second equation of fluid dynamics:

$$\frac{\partial \mathbf{v}}{\partial t} + (\mathbf{v} \cdot \nabla) \mathbf{v} = -\frac{\nabla P}{\rho} - \nabla \Phi \quad (7.3)$$

where P is the pressure and Φ is the gravitational potential. We can write this equation in terms of Lagrangian derivatives as:

$$\frac{D\mathbf{v}}{Dt} = \mathbf{g} - \frac{\nabla P}{\rho} \quad (7.4)$$

In the absence of external or pressure forces, the fluid velocity will remain unchanged in the Lagrangian frame.

There is a third key equation in fluid dynamics – energy conservation – but we will not require that in the arguments below so do not discuss it here.

We can now consider the generalization to phase space. We define the **distribution function (DF)**, $\mathbf{f}$, as a generalization of density in 6-dimensional phase space (x, y, z, v_x, v_y, v_z) :

$$\int f(x, v, t) d^3x d^3v = 1 \quad (7.5)$$

The DF is simply the probability of finding a body (star, etc.) in a 6D phase space element $d^3x d^3v$.

Many expressions connected to observations can be obtained by taking moments (i.e., integrals) of the DF.

For example, we can obtain the density via:

$$n(\mathbf{x}, t) = N \int f d^3v = N\nu \quad (7.6)$$

where N is the total number of stars in the system. Note that we will write density here as n , i.e., the *number* density, and ν is the normalized number density, i.e., the probability in configuration space.

The mean velocity can be obtained via:

$$\bar{\mathbf{v}} = \frac{\int \mathbf{v} f d^3v}{\int f d^3v} \quad (7.7)$$

and so on.

We are now going to derive the fundamental equation governing the evolution of the DF. We begin by defining $\mathbf{w} = (\mathbf{x}, \mathbf{v})$. Note that this means we could have written our earlier equation as $\int f(w) d^6w = 1$, or $\int f(\mathbf{w}) d\mathbf{w} = 1$. We then have:

$$\dot{\mathbf{w}} = (\dot{\mathbf{x}}, \dot{\mathbf{v}}) = (\mathbf{v}, -\nabla\Phi) \quad (7.8)$$

By analogy with the fluid equations, conservation of probability directly entails:

$$\frac{\partial f}{\partial t} + \frac{\partial}{\partial \mathbf{w}} \cdot (f \dot{\mathbf{w}}) = 0 \quad (7.9)$$

we can expand the second term so that the equation becomes:

$$\frac{\partial f}{\partial t} + f \frac{\partial \dot{\mathbf{w}}}{\partial \mathbf{w}} + \dot{\mathbf{w}} \frac{\partial f}{\partial \mathbf{w}} = 0 \quad (7.10)$$

We can further expand the second term as:

$$\frac{\partial \dot{\mathbf{w}}}{\partial \mathbf{w}} = \frac{\partial \mathbf{v}}{\partial \mathbf{x}} + \frac{\partial(-\nabla\Phi)}{\partial \mathbf{v}} \quad (7.11)$$

both of the terms on the right hand side are zero: velocities and positions are not correlated, and the potential does not depend on velocity. This allows us to write the conservation of probability as the more familiar **collisionless Boltzmann equation** (CBE; also known as the Vlasov equation):

$$\frac{\partial f}{\partial t} + \mathbf{v} \cdot \nabla f - \nabla\Phi \cdot \frac{\partial f}{\partial \mathbf{v}} = 0 \quad (7.12)$$

If we define:

$$\frac{Df}{Dt} \equiv \frac{\partial f}{\partial t} + \dot{\mathbf{w}} \frac{\partial f}{\partial \mathbf{w}} \quad (7.13)$$

then we can again write the CBE in Lagrangian coordinates, in which case the equation becomes much simpler:

$$\boxed{\frac{Df}{Dt} = 0} \quad (7.14)$$

One way to interpret this equation is that the DF obeys incompressible flow in phase space. Phrased another way, the DF is a conserved quantity along orbits. We will see in a later chapter that this fact can be used to place constraints on how galaxies have evolved over cosmic time. For example, if two galaxies merge (and do not form new stars) then the phase space densities will be the same pre and post merger.

Notice that Poisson's equation was not used to derive the CBE. This means that the particles obeying the CBE need not be the same particles generating the potential. In other words, the particles obeying the CBE could be *tracers* of the underlying mass distribution. This is important: we can use stars, which will satisfy the CBE, as tracers of the underlying *mass* distribution. The most common type of tracer that we will encounter in this class is stars.

Note that the CBE does not consider stellar evolution (stellar death), nor the formation of new stars. It also does not include or consider correlations between stars (either because stars form in a correlated way or because gravity induces correlations). These are all limitations of the CBE, although one can formulate extensions that take these complications into account.

One can fill many pages (BT certainly does) with examples of DFs, general theorems, alternative formulations, etc. Here we will focus on just a few to give a flavor of the field.

A DF that is a function only of $E = \frac{1}{2}v^2 + \Phi$, i.e., $f = f(E)$ is said to be *ergodic*, or *isotropic*. Such systems are also spherically symmetric. In the case of an ergodic DF, the normalized density profile, $\nu(r)$, and the spherically symmetric (but not necessarily self-gravitating) potential $\Phi(r)$ produces a unique $f(E)$ that is given by Eddington's formula:

$$f(E) = \frac{1}{\sqrt{8\pi^2}} \frac{d}{dE} \int_E^\infty \frac{d\nu}{d\Phi} \frac{d\Phi}{\sqrt{\Phi - E}} \quad (7.15)$$

There is a class of DFs that obey $f(E) \propto E^k$ that are known as polytropes in which the polytropic index n is related to k via $k = n - 3/2$. The case of $n = 5$ is the Plummer model that we introduced in the previous chapter. Note that polytropes are assumed to be self-gravitating, so that Poisson's equation can be used to relate ρ and Φ .

The isothermal sphere is a special case of the polytrope, in which $n \rightarrow \infty$. In this case the DF is:

$$f(E) = K e^{-E/\sigma^2} = K e^{-\Phi/\sigma^2} e^{-v^2/2\sigma^2} \quad (7.16)$$

notice that the second term in the second expression is the usual Maxwell-Boltzmann distribution of velocities for a thermal distribution. The solution is not analytic except for when $\Phi = V_c^2 \ln r$. In this case:

$$n(r) = \frac{\sigma^2}{2\pi G r^2} \quad (7.17)$$

and $V_c^2 = 2\sigma^2$.

I'll leave it as an exercise to the reader to show that the mean-squared speed:

$$\overline{v^2} = \frac{\int f(E) v^2 d^3v}{\int f(E) d^3v} = 3\sigma^2 \quad (7.18)$$

where $\sigma = \text{constant}$. i.e., an isothermal model has a constant velocity dispersion profile (we can associate σ^2 with a “temperature”, hence the language of isothermal).

Beyond ergodic DFs, the next level of complexity is to consider models of the form $f(E, L^2)$, i.e., anisotropic DFs. Because L^2 is independent of direction, these DFs are still spherically symmetric. However, the symmetry between radial and tangential directions is broken, and we can now have more stars on radial or circular (tangential) orbits, or vice versa. In such models the **anisotropy parameter**, β , is introduced:

$$\beta \equiv 1 - \frac{\sigma_\theta^2 + \sigma_\phi^2}{2\sigma_r^2} = 1 - \frac{\sigma_\theta^2}{\sigma_r^2} \quad (7.19)$$

where we have assumed symmetry between the θ and ϕ components in the second equality. Models in which $\beta = 0$ are isotropic; models with $\beta = 1$ are dominated by radial orbits, while models with $\beta \rightarrow -\infty$ are dominated by circular orbits. This might seem like a detail, but as we will see below, knowledge of β is essential for determining the masses of dynamical systems.

A further refinement is to consider axisymmetric DFs of the form $f(E, L_z)$; these are important for modeling the dynamics of disk-dominated galaxies.

In summary, we can say that the DF and Φ determine the distribution of particles (stars) in x and v , e.g., the density distribution and velocity distribution. The CBE governs the subsequent temporal evolution. These relations are very powerful, but analytic examples are scarce.

The concept of **phase mixing** plays an important role in galactic dynamics. As a reminder, the distribution function (DF), f , obeys the collisionless Boltzmann equation and is a *conserved* quantity. This is sometimes referred to as the *fine-grained* DF, to contrast with the *coarse-grained* DF. The latter, often denoted $\bar{f}$, is not conserved, and is an average of the DF

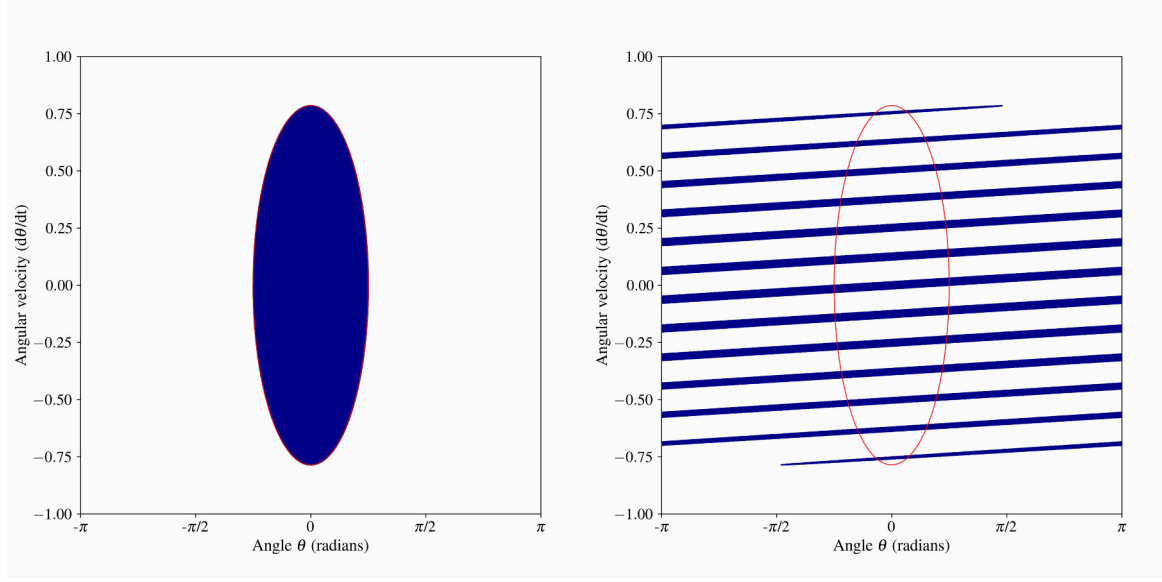


Figure 7.1: A toy model illustrating the concept of phase mixing. A region in the 2D phase space of $\theta - \dot{\theta}$ is initialized at left. After a period of time (and evolved without any forces), the region is sheared out over a large swath of phase space. At a microscopic level the distribution function f is conserved. However, the coarse-grained function, $\bar{f}$ is strictly $< f$. A video of this simulation is available [here](#).

over a finite volume in phase space. Observationally, we can only ever measure the coarse-grained DF, and so we can expect $\bar{f}$ to decrease with time. This is known as phase mixing. Essentially what is happening is that by averaging over a finite volume, we will in general be including regions of empty phase space. See BT Ch 4.10.2 and Fig. 7.1 for examples. Note that phase mixing occurs even when the energies of individual particles are conserved.

We can gain additional insight by taking velocity moments of the CBE. The 0th moment, $\int \text{CBE } d\mathbf{v}$, just returns the continuity equation, i.e., the conservation of “mass” for the tracer particles:

$$\frac{\partial \nu}{\partial t} + \nabla \cdot (\nu \mathbf{v}) = 0 \quad (7.20)$$

The first moment is more interesting, though it’s rather complicated to derive (see BT Ch 4.8 for details). In spherical systems the first moment simplifies to:

$$\frac{d(\nu \overline{v_r^2})}{dr} + 2 \frac{\beta}{r} \nu \overline{v_r^2} = -\nu \frac{d\Phi}{dr} \quad (7.21)$$

These two equations are known as the **Jeans equations**.

Recall again that ν need not be the same mass that generates the potential (i.e., it can be a tracer population), and likewise $\overline{v_r^2}$ is the mean squared radial velocity of the tracers. Therefore, leaving β aside, the above equation has two observables, and one unknown - the potential. *The Jeans equation is our route to determine the gravitating mass (total mass) of a system based on a tracer population.*

The total mass interior to radius r can be derived from the Jeans equations and expressed in a revealing way:

$$M_{\text{tot}}(r) = \frac{r\sigma_r^2}{G} (\gamma_* + \gamma_\sigma - 2\beta) \quad (7.22)$$

where $\gamma_* \equiv -d \ln \nu / d \ln r$ and $\gamma_\sigma \equiv -d \ln \sigma_r^2 / d \ln r$, i.e., the logarithmic slopes of the tracer density profile and radial velocity dispersion profile. Expressed this way, one sees clearly the importance of the anisotropy parameter in the mass determination.

One of the central challenges in using dynamics to infer the mass distribution is setting, or measuring the anisotropy profile. Often $\beta(r)$ is simply taken from simulations (which is not great practice). In principle, β can be measured from the data itself, as the higher-order moments of the velocity distribution contain information on the anisotropy. But in practice this is very difficult. Further, keep in mind that the anisotropy enters twice in the mass modeling: both directly (as seen above) and because the *radial* velocity dispersion is never measured; instead we measure the line-of-sight velocities, which are determined by both σ_r^2 and β (as well as $\nu(r)$).

A very useful model is the **Jeans equation for a 1D slab**:

$$\frac{\partial(\nu\sigma_z^2)}{\partial z} = -\nu \frac{\partial\Phi}{\partial z} \quad (7.23)$$

We can write Poisson's equation in 1D as:

$$4\pi G\rho = \frac{\partial^2\Phi}{\partial z^2} \quad (7.24)$$

where remember that ρ is the mass density producing the potential and need not be the same as ν .

Combining these two equations yields:

$$\rho(z) = \frac{-1}{4\pi G} \frac{\partial}{\partial z} \left[\frac{1}{\nu} \frac{\partial(\nu\sigma_z^2)}{\partial z} \right] \quad (7.25)$$

In principle we can measure the density in the z direction from measurements of ν and σ_z – both observables. The challenge in this case is that the measurement requires double differentiation of a density (triple differentiation of star counts)– in general that is going to be a very noisy measurement.

This idealization of a 1D slab is a reasonably good approximation for the mass distribution in the disk of our Galaxy (and others), where the disk scale height is small compared to the radial scale length (so that the radial derivatives can be neglected in the full Jeans equation). Use of the 1D Jeans equation is the origin of the **Oort Limit** on the total mass density in the solar neighborhood. Current limits place the density at the mid-plane of $\rho_{\text{dyn}}(z = 0) \approx 0.08 - 0.10 M_{\odot} \text{pc}^{-3}$. A careful accounting of all the baryonic matter in the solar neighborhood (including stars, gas, and stellar remnants) gives a value of $\rho_{\text{bary}} \approx 0.09 M_{\odot} \text{pc}^{-3}$; the conclusion is that there is no evidence of extra mass within the solar neighborhood. This is not in tension with modern cosmological models as most predict a small local density of dark matter.

Chapter 8

The Virial Theorem & Dynamical Friction

In this chapter we continue our detour into the world of dynamics by introducing the virial theorem and the concept of dynamical friction. The virial theorem is very widely used in astronomy, often as a “back of the envelope” way to estimate masses given velocities and sizes. As we will see, the virial theorem was employed to obtain the very first evidence for dark matter. Dynamical friction is a complex, collective dynamical process that is important for the dynamical evolution of many systems in astronomy.

The virial theorem can be derived in several ways. A very thorough version can be found in BT Ch 4.8.3 and involves integrating the Jeans equation multiplied by $\mathbf{x}$. The classic derivation involves defining the quantity:

$$G \equiv \sum_i \mathbf{p}_i \cdot \mathbf{r}_i \quad (8.1)$$

where $\mathbf{p}$ and $\mathbf{r}$ are the momenta and positions of particles. Setting $\mathbf{p} = m\mathbf{v} = m\dot{\mathbf{r}}$ allows us to write:

$$G = \sum_i m_i \frac{d\mathbf{r}_i}{dt} \cdot \mathbf{r}_i = \frac{1}{2} \sum_i m_i \frac{d(\mathbf{r}_i \cdot \mathbf{r}_i)}{dt} = \frac{1}{2} \frac{d}{dt} \left(\sum_i m_i r_i^2 \right) \quad (8.2)$$

where the last term in brackets is just the moment of inertia, I . We therefore have:

$$G = \frac{1}{2} \frac{dI}{dt} \quad (8.3)$$

The next step is to take the time derivative of G :

$$\frac{dG}{dt} = \sum_i \dot{\mathbf{p}}_i \cdot \mathbf{r}_i + \sum_i \mathbf{p}_i \cdot \dot{\mathbf{r}}_i \quad (8.4)$$

where $\mathbf{p}_i$ is just the applied force, $\mathbf{f}_i$, and

$$\sum_i \mathbf{p}_i \cdot \dot{\mathbf{r}}_i = \sum_i m_i \dot{\mathbf{r}}_i \cdot \dot{\mathbf{r}}_i = \sum_i m_i v_i^2 = 2K \quad (8.5)$$

where K is the total kinetic energy of the system. It takes a bit more work to show that the term $\sum_i \mathbf{f}_i \cdot \mathbf{r}_i$ is equal to the total potential energy of the system, W . Collecting what we have learned:

$$\frac{dG}{dt} = \frac{1}{2} \frac{d^2 I}{dt^2} = 2K + W \quad (8.6)$$

If we now set $\frac{dG}{dt} = 0$ then we arrive at the **scalar virial theorem**:

$$\boxed{W + 2K = 0} \quad (8.7)$$

A few comments are in order. First, the assumption of $\frac{dG}{dt} = 0$ is equivalent to $\ddot{I} = 0$. Thus, the virial theorem holds if the system is stationary. Notice that we did not assume any time averaging, and so the usual virial theorem can be applied to any system in which $\ddot{I} = 0$ at any moment in time. One can obtain alternative versions of the virial theorem in which the quantities are all taken to be some temporal average, in which case the system need not be instantaneously stationary, but only stationary over some timescale, i.e., $\langle \ddot{I} \rangle = 0$. This form is less useful in astronomy as the relevant timescales are often millions or billions of years.

Notice that we assumed $\ddot{I} = 0$ but not $\dot{I} = 0$. If $\dot{I} \neq 0$ then there are net radial flows of momentum within the system. Technically the virial theorem would still hold in this case, but such situations are highly contrived. In all cases that we will encounter, when the system is in equilibrium we will have both $\ddot{I} = 0$ and $\dot{I} = 0$.

We also know that the total energy, E , is $E = K + W$ and so we can write various versions of the virial theorem, all of which are equivalent: $W = -2K$, $E = -K$, $E = \frac{1}{2}W$, etc.

A complete derivation of the virial theorem includes a surface pressure term, Σ :

$$\frac{1}{2} \ddot{I} = W + 2K + \Sigma \quad ; \quad \Sigma = - \int_S \rho \sigma^2 \mathbf{x} \cdot d\mathbf{S} \quad (8.8)$$

In many astrophysical contexts the surface pressure is either zero or negligible compared to the other terms, so it can safely be ignored.

The virial theorem reveals an important feature of self-gravitating systems. Consider a system of particles at rest at infinity. In this initial configuration we have $E_i = W_i = K_i = 0$ for the system. Now collect the material into an equilibrium configuration ($\ddot{I} = 0$) so that the virial theorem is satisfied. The final configuration then obeys: $E_f = -K_f = \frac{1}{2}W_f$. Considering that energy must be globally conserved, this tells us that half of the gravitational

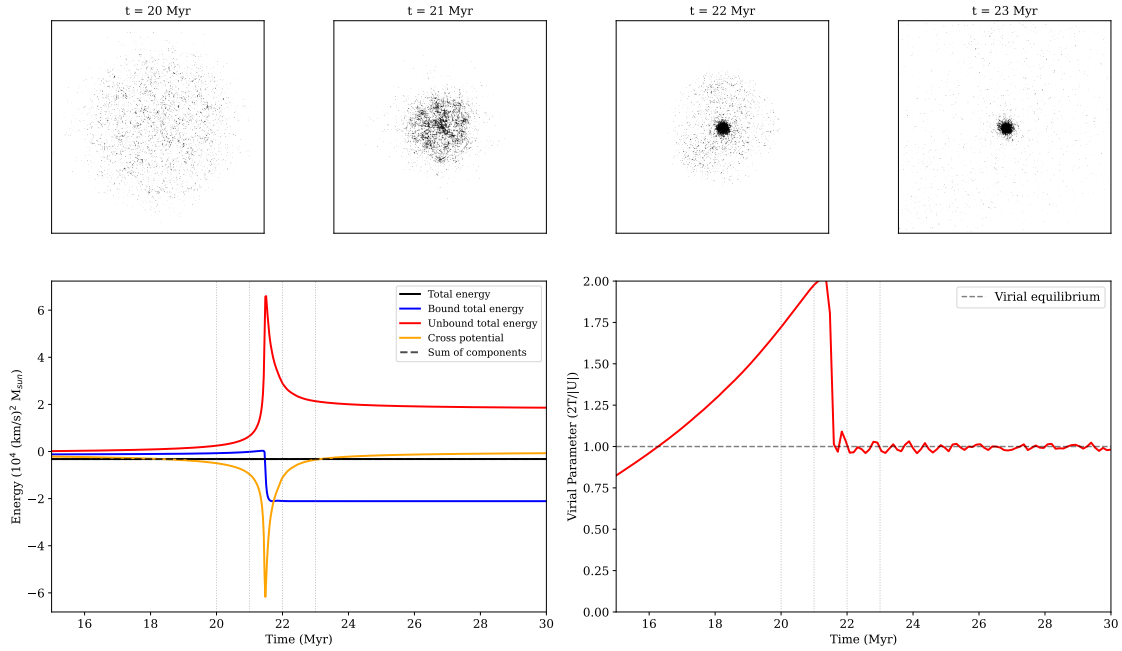


Figure 8.1: Collapse of a system of particles initially at rest with very little total energy. The top panels show four snapshots of the system around the moment of collapse and virialization. The bottom panels show the energy of various components and the virial parameter for the bound particles as a function of time. A video of this simulation is available [here](#).

energy goes into kinetic form, and the other half is “disposed of”. But where does the other half of the energy go? There are a variety of mechanisms available. For example, in the collapse of a gas cloud, the energy is radiated away. In the collapse of a collisionless system (e.g., stars, or dark matter), the energy is carried away by ejected particles (see Fig. 8.1). One can phrase the situation differently by saying that a system *cannot* reach dynamical equilibrium *until* it has shed half of its final binding energy.

The most common use of the virial theorem is to relate the mass, velocity dispersion, and size of a system. This is done by setting $K = \frac{1}{2}M\sigma^2$ and $W = -GM^2/(\alpha R)$, where $1 \lesssim \alpha \lesssim 5$ is a dimensionless parameter that captures the mass distribution of the system. The virial theorem then implies:

$$M = \alpha \frac{\sigma^2 R}{G} \quad (8.9)$$

You might think that this very simple approximation is too crude to be useful. But you would be wrong.

In 1933 Fritz Zwicky published results on a spectroscopic survey of galaxies in the Coma cluster (the most nearby truly massive cluster). The dispersion in redshifts was measured to be $\sigma^2 = 5 \times 10^{15} \text{ cm}^2 \text{ s}^{-2}$, and the size of the cluster was estimated at $R \approx 1 \text{ Mpc}$. From

these observations Zwicky used the virial theorem to deduce $M_{\text{tot}} \approx 10^{14} M_{\odot}$. There are $\sim 10^3$ galaxies in Coma, each with $L \sim 8 \times 10^7 L_{\odot}$, resulting in an average mass-to-light ratio of $M/L \sim 1000$ (units are rarely explicitly noted for M/L , but they are assumed to be $M_{\odot}/L_{\odot}$). Compare this to $M/L \sim 3$ in the solar neighborhood. This is a remarkable discrepancy, and Zwicky listed the possible explanations: 1) stars within Coma galaxies are very different from the solar vicinity (remember that this was 1933, and the basic theory of stellar evolution was still in its infancy). 2) Coma is not in dynamical equilibrium (and therefore the virial theorem is not applicable). 3) The laws of physics are not the same in the Coma cluster (or, phrased in modern parlance, the laws of physics at the *scales* relevant for the measurement are not the same as the laws measured and tested in the solar system). 4) There exists large quantities of invisible (dark) matter. The debate between options 3 and 4 (MOND vs. dark matter) is still alive and well today, although at present the preponderance of the data favors option 4.

This example offers an important general lesson: sometimes in science very careful, detailed calculations are necessary to arrive at the correct answer (consider analysis of the CMB acoustic oscillations), but in other cases “factor-of-two” analysis is sufficient to provide fundamental insight. One of the characteristics of a successful scientist is to know when to apply which technique to a given problem.

More recently, sophisticated Jeans-based dynamical models of galaxies have been compared to simple virial estimates; the agreement is remarkably good with a scatter of only 15% (see Fig. 8.2).

We are now going to briefly discuss the thermodynamics of self-gravitating systems with the aid of the virial theorem. Let’s work toward defining the heat capacity, i.e., the amount of energy required to change the temperature by one degree. By analogy with an ideal gas, we can define a temperature via $\frac{1}{2}m\langle v^2 \rangle = \frac{3}{2}kT$ where k is Boltzmann’s constant. In general we would want to consider some appropriate average temperature, e.g., a mass-weighted average: $\langle T \rangle = \int d^3x \rho T / \int d^3x \rho$. The total kinetic energy of the system can then be expressed as: $K = \frac{3}{2}Nk\langle T \rangle$. The virial theorem tells us that $E = -K = -\frac{3}{2}Nk\langle T \rangle$. The heat capacity is then:

$$C \equiv \frac{dE}{d\langle T \rangle} = -\frac{3}{2}Nk \quad (8.10)$$

the key point is that the *heat capacity is negative*, which is very unusual if one is used to thinking about ideal gases. By losing energy a system grows hotter. This result is very general: any bound, finite system dominated by gravity has a negative heat capacity.

Any system with a negative heat capacity is thermally unstable. Imagine putting a self-

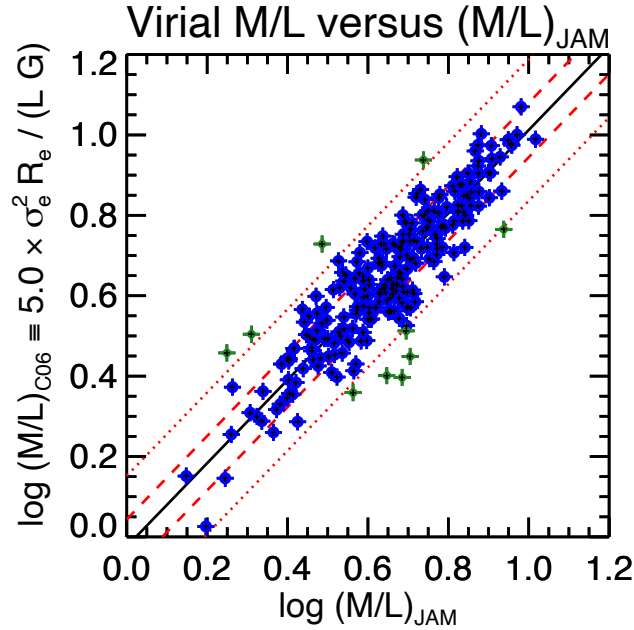


Figure 8.2: Comparison of virial and Jeans-based dynamical masses for elliptical galaxies from the ATLAS^{3D} sample. From Cappellari et al. (2013).

gravitating fluid in contact with a heat bath, initially at the same temperature as the bath. Imagine transferring a small amount of energy from the fluid to the bath. Since $C < 0$, the fluid heats up (by definition the bath does not). Heat will then continue to flow from hot to cold (from the fluid to the bath), resulting in runaway increase in the temperature of the fluid. This is the origin of “core collapse” in globular clusters.

Dynamical friction is the process by which kinetic energy is drained from an orbiting satellite, eventually leading to the merging of the satellite with the host system. Aspects of the derivation below will echo our earlier analysis of two-body relaxation. Below is a relatively detailed sketch; the proper derivation can be found in BT Ch 7.4.4 and Ch 8.1. Dynamical friction was first worked out in detail by Chandrasekhar (1943) and so the formula often bears his name.

Consider a massive body of mass M moving with velocity v through a homogeneous sea of particles at rest, each of mass m , such that $M \gg m$, and with a number density of the background sea of n and mass density $\rho = nm$. We will begin by considering the change in velocity of a single small particle due to its interaction with the large body (for this calculation we move into the frame of the large body and so the small body is moving with velocity v relative to the large body). We have $\Delta v_{\parallel} = 0$ by symmetry. In the perpendicular

direction we have:

$$\Delta v_{\perp} = \int a_{\perp} dt = \int_{-\infty}^{+\infty} dt \frac{GM}{r^2} \frac{b}{r} \quad (8.11)$$

where b is the impact parameter of the encounter, and r is the separation between the two bodies. Setting $z = vt$, where z is the distance traveled along the direction of motion, we have:

$$\Delta v_{\perp} = \frac{GM}{v} \int_{-\infty}^{+\infty} dz \frac{b}{(b^2 + z^2)^{3/2}} = \frac{2GM}{bv} \quad (8.12)$$

the change in kinetic energy of the small particle is then:

$$\Delta E = \frac{1}{2} m \Delta v^2 = \frac{2G^2 M^2 m}{b^2 v^2} \quad (8.13)$$

The number of encounters between b and $b + db$ over a time interval dt is $2\pi b db n v dt$. The background sea of stars therefore acquires energy at the rate:

$$\frac{dE}{dt} = 2\pi \int_0^{\infty} b db n \Delta E v \quad (8.14)$$

$$= \frac{4\pi G^2 M^2 n m \ln \Lambda}{v} \quad (8.15)$$

where $\ln \Lambda$ is the Coulomb logarithm. Energy gained by the background sea is equal to the energy lost by the large body:

$$-\frac{dE}{dt} = -\frac{d(\frac{1}{2} M v^2)}{dt} = -M v \frac{dv}{dt} \quad (8.16)$$

so:

$$\frac{dv}{dt} = -\frac{4\pi G^2 M \rho \ln \Lambda}{v^2} \quad (8.17)$$

Recall that $\Lambda \equiv b_{\max}/b_{\min}$, and we can approximate $b_{\max} \approx R$, where R is the size of the system (i.e., of the background sea), and $b_{\min} = \max(GM/v^2, R_M)$ where R_M is the size of the large body.

We can define the **dynamical friction timescale**, t_{DF} as the time it takes for the large body's initial velocity to be reduced to zero:

$$\frac{dv}{dt} = \frac{(0 - v)}{t_{\text{DF}}} \quad (8.18)$$

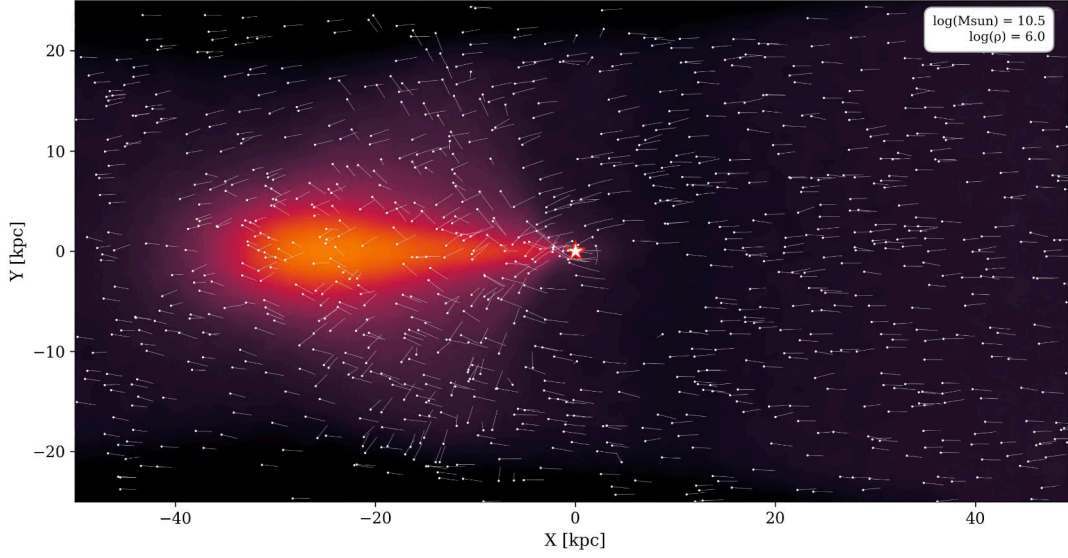


Figure 8.3: Dynamical friction wake trailing a massive object. The massive body is at center, with a mass of $3 \times 10^{10} M_{\odot}$ distributed in a Plummer profile. The background is a sea of particles with a combined mass density of $\rho = 10^6 M_{\odot} \text{ kpc}^{-3}$. The relative velocity between the massive object and the background is 200 km s^{-1} . Test particles (white) show the direction of the flow. A video of this simulation is available [here](#).

which leads us to:

$$t_{\text{DF}} = \frac{v^3}{4\pi G^2 M \rho \ln \Lambda} \quad (8.19)$$

The dynamical friction timescale is inversely proportional to both the background density, ρ , and the large body mass, M . Both of these are intuitive – with more background sea particles the encounter rate will be larger, and each individual encounter imparts a velocity shift proportional to the large body mass. The timescale is also proportional to the third power of velocity. The rate of energy gain from the background per encounter goes as v^{-2} , but the encounter *rate* goes as v , so the total energy gain rate goes as v^{-1} . The conversion from energy rate to velocity rate adds another factor of v^{-1} , and the final power comes from the fact that a larger velocity means it will take longer to reduce the velocity to zero.

The derivation above was idealized in several respects: the background density was constant, as was the mass and velocity of the large body. In reality, for any realistic orbit the body will be moving through a background of varying density. Other processes such as tidal stripping will result in a reduction of the large body mass with time. Dynamical friction also heats the background sea, resulting in a changing velocity distribution over time. All of these effects are complex, and most modern quantitative approaches rely on numerical simulation. See

Fig. 8.3.

The same basic physical process will occur when the background is a gas, with correction factors appearing when $v/c_s \gtrsim 1$, where c_s is the sound speed.

Another way to understand dynamical friction is as a manifestation of **equipartition**. The basic idea is that two-body encounters move systems toward equilibrium in which $m_1 v_1^2 = m_2 v_2^2$. In the case of dynamical friction we have $M \gg m$, and in many situations of interest $v^2 \approx \sigma^2$ (the velocity of the body is comparable to the typical velocity of the particles), so the tendency toward equipartition results in the larger body slowing down relative to the background sea.

Dynamical friction drains kinetic energy from a body moving in a sea of smaller bodies and is the primary mechanism that results in the decay of orbits in galactic contexts. Let's consider some examples:

- Imagine a supermassive black hole in orbit around a galaxy with an isothermal density profile: $\rho(r) = \frac{\sigma^2}{2\pi G r^2}$. The dynamical friction timescale is:

$$t_{\text{DF}} = \frac{19 \text{ Gyr}}{\ln \Lambda} \left(\frac{R_i}{5 \text{ kpc}} \right)^2 \left(\frac{\sigma}{200 \text{ km s}^{-1}} \right) \left(\frac{10^8 M_\odot}{M_{\text{BH}}} \right) \quad (8.20)$$

where R_i is the initial orbital radius of the BH. Assuming $\ln \Lambda = 6$ we have $t_{\text{DF}} = 3 \text{ Gyr}$ for typical parameters. Supermassive BHs will sink to the center of a galaxy within a Hubble time. Notice in the above equation that the apparent velocity dependence is only to the first power – this is because the density profile scales as v^2 and the orbital velocity of the BH is of course ultimately set by the potential.

Recall that $M(r) \sim v^2 r / G$ and $t_{\text{cross}} = r/v$, so we can recast the above equation as:

$$t_{\text{DF}} = \frac{1.2}{\ln \Lambda} \frac{M(r)}{M_{\text{BH}}} t_{\text{cross}} \quad (8.21)$$

In words, the dynamical friction timescale is a multiple of the crossing time, where the factor is determined by the ratio of enclosed mass to the object mass.

- Imagine a globular cluster (GC) orbiting in the same potential as above. Most Galactic GCs have masses of $\sim 2 \times 10^5 M_\odot$. Adopting this as the fiducial GC mass, we have:

$$t_{\text{DF}} = 64 \text{ Gyr} \left(\frac{R_i}{1 \text{ kpc}} \right)^2 \left(\frac{\sigma}{200 \text{ km s}^{-1}} \right) \quad (8.22)$$

where we've assumed $\ln\Lambda = 5.8$ ($b_{\max} = 1$ kpc and $b_{\min} = 3$ pc where the latter is the typical size of a GC).

GCs will spiral into the centers of galaxies if their initial radius is small, or if the GCs are in orbit around small galaxies (e.g., with $\sigma \sim 10$ km s⁻¹).

- In galaxy clusters the typical velocity is $\sim 10^3$ km s⁻¹ and sizes are of order 1 Mpc. In clusters the dynamical friction timescale is much longer than the age of the universe. Galaxy-galaxy mergers are rare in clusters; this is one reason why clusters contain many satellite galaxies. Galaxy-galaxy mergers are much more common in the galaxy group scale, where velocities are of order a few hundred km s⁻¹.
- As we will see in later chapters, dark matter halos have a fixed, constant density within their virial radius by definition, i.e., $\rho(R_{\text{vir}}) = \text{constant}$. This in turn implies that the total halo mass, M_h , is proportional to v^3 . The dynamical friction time of a smaller satellite dark matter halo orbiting within a larger host dark matter halo is then:

$$t_{\text{DF}} \approx \frac{0.1}{\ln(1 + M_{\text{host}}/M_{\text{sat}})} \frac{M_{\text{host}}}{M_{\text{sat}}} t_H \quad (8.23)$$

where the Coulomb logarithm was estimated from N -body simulations. Satellite halos can merge with the central host within a Hubble time; the key variable is the *ratio* of host and satellite masses.

What follows is a back-of-the-envelope calculation for the dynamical friction timescale that provides another view into the physical processes at work.

The gravitational deflection of the streamlines caused by the massive body results in an enhancement of mass behind the body. Call this mass enhancement Δm . Let R be the distance over which the streamlines are diverted by the body. On dimensional grounds we can approximate $R \sim GM/v^2$. Let $\dot{m} = \pi R^2 \rho v$ be the mass flux of the background sea of particles through the affected area. Then $\Delta m = \dot{m} \Delta t \sim \dot{m} R/v$. The force exerted *by* the mass enhancement *on* the body is:

$$F = \frac{GM\Delta m}{R^2} = M\dot{v} = M \frac{v}{t_{\text{DF}}} \quad (8.24)$$

re-arranging:

$$t_{\text{DF}} \sim \frac{vR^2}{G\Delta m} \sim \frac{vR^2}{G\pi R^2 \rho v R/v} \sim \frac{v}{GR\rho} \sim \frac{v^3}{G^2 M \rho} \quad (8.25)$$

Compare to Eq. 8.19.

Part IV

Growth of Structure

Chapter 9

Linear Fluctuations

Having completed a tour of galactic dynamics, we now return to the story of cosmology. In earlier chapters we presented the governing equations for the evolution and structure of a homogeneous universe, i.e., the behavior of the universe on very large ($\gtrsim 100$ Mpc) scales. Our task now is to explain how tiny fluctuations in the density field, presumably seeded by quantum fluctuations, grow into the spectacular non-linear structures we observe today: galaxies, clusters, and other aspects of **large scale structure** in the universe.

In this chapter and the next we will focus on the linear growth of perturbations.

A useful way to quantify structure in the universe is the dimensionless density contrast:

$$\delta \equiv \frac{\rho - \rho_0}{\rho_0} = \frac{\rho - \bar{\rho}}{\bar{\rho}} = \frac{\delta\rho}{\rho} \quad (9.1)$$

An advantage of δ is that it is dimensionless, although a notable disadvantage is that it is bounded on one side: $-1 \leq \delta < \infty$. At the present epoch, we can quantify the approximate density contrast of several aspects of large-scale structure. The intergalactic medium (IGM) has $1 \lesssim \delta \lesssim 10$ in filaments; the largest superclusters have $\delta \sim 2 - 5$ when averaged over ~ 10 Mpc; dark matter halos have $\delta \approx 200$ (partly by definition); galaxies have $\delta \sim 10^6$.

By linear fluctuations we mean $\delta \ll 1$, i.e., the density fluctuations are small and, as we will see below, this allows us to neglect terms beyond first-order. We will see in a later chapter that once $\delta \sim 1$ the subsequent collapse into a bound object occurs quite rapidly.

We are going to begin by considering the behavior of linear fluctuations within a self-gravitating fluid in a static (non-expanding) universe. We will employ the fluid equations

(continuity and Euler) in Lagrangian form as well as Poisson's equation:

$$\frac{d\rho}{dt} = -\rho \nabla \cdot \mathbf{v} \quad (9.2)$$

$$\frac{dv}{dt} = -\frac{1}{\rho} \nabla P - \nabla \Phi \quad (9.3)$$

$$\nabla^2 \Phi = 4\pi G \rho \quad (9.4)$$

We will consider small fluctuations on top of a uniform, zero-velocity background density, so that: $\mathbf{v} = \delta \mathbf{v}$; $\rho = \rho_0 + \delta \rho$; $P = P_0 + \delta P$; $\Phi = \Phi_0 + \delta \Phi$, where the background is given by ρ_0 , P_0 and Φ_0 , and δx are the perturbed quantities¹. Note that our setup contains a serious flaw: it is not possible to have a uniform distribution of matter sitting at rest. This is known as the Jeans swindle, and there are various attempts to provide some justification for it (see e.g., BT Ch 5.2.5). For our purposes we will adopt the Jeans swindle and move on.

Our aim is to insert the perturbed quantities into the fluid equations and keep only terms to first order. For the continuity equation we have:

$$\frac{d(\rho_0 + \delta \rho)}{dt} = -(\rho_0 + \delta \rho) \nabla \cdot \delta \mathbf{v} \quad (9.5)$$

we can ignore the second term on the RHS as that is of second-order. We can then subtract the unperturbed continuity equation ($d\rho_0/dt = -\rho_0 \nabla \cdot \mathbf{v}_0$, where recall that $\mathbf{v}_0 = 0$) to arrive at our first key equation:

$$\frac{d\delta}{dt} = -\nabla \cdot \delta \mathbf{v} \quad (9.6)$$

For the second fluid equation, we have:

$$\frac{d\delta \mathbf{v}}{dt} = -\frac{\nabla(P_0 + \delta P)}{\rho_0 + \delta \rho} - \nabla(\Phi_0 + \delta \Phi) \quad (9.7)$$

we have assumed a uniform background, so $\nabla P_0 = \nabla \Phi_0 = 0$. Keeping only terms at first order results in the second key equation:

$$\frac{d\delta \mathbf{v}}{dt} = -\frac{1}{\rho_0} \nabla \delta P - \nabla \delta \Phi \quad (9.8)$$

Finally, because of the linearity of Poisson's equation, we have for the perturbed quantities:

$$\nabla^2 \delta \Phi = 4\pi G \delta \rho \quad (9.9)$$

¹Be aware of the somewhat confusing notation: in general δx refers to a small perturbation to a quantity x ; while δ on its own refers to the density contrast.

Taking a time derivative of Eq. 9.6, we have:

$$\ddot{\delta} = -\nabla \cdot \dot{\delta \mathbf{v}} \quad (9.10)$$

and taking the divergence of Eq. 9.8 we have:

$$\nabla \cdot \dot{\delta \mathbf{v}} = -\frac{1}{\rho_0} \nabla^2 \delta P - \nabla^2 \delta \Phi \quad (9.11)$$

We will assume that the perturbations are adiabatic, which means that we can write $c_s^2 = \partial P / \partial \rho$ where c_s is the sound speed. This allows us to set $\delta P = c_s^2 \delta \rho$. Finally, combining Eq. 9.9, 9.10, and 9.11 we arrive at:

$$\boxed{\ddot{\delta} = c_s^2 \nabla^2 \delta + 4\pi G \rho_0 \delta} \quad (9.12)$$

This is the governing equation for the evolution of linear adiabatic self-gravitating fluctuations.

The final step is to assume a wave solution to the above equation. We take the usual approach and assume $\delta \propto e^{i(\mathbf{k} \cdot \mathbf{x} - \omega t)}$. Inserting this into Eq. 9.12 we arrive at the dispersion relation:

$$\omega^2 = c_s^2 k^2 - 4\pi G \rho_0 \quad (9.13)$$

Now for the interpretation: for $c_s^2 k^2 > 4\pi G \rho_0$, the perturbations are oscillatory, i.e., they are sound waves supported by the pressure gradient. For $c_s^2 k^2 < 4\pi G \rho_0$ the frequencies are imaginary, which means the modes are unstable. The over-density cannot be stabilized by the pressure gradient, and collapse will ensue. The critical scale is:

$$\lambda_J = \frac{2\pi}{k_J} = c_s \left(\frac{\pi}{G \rho_0} \right)^{1/2} \quad (9.14)$$

This is called the **Jeans length**, and there is a corresponding Jeans mass: $M_J = \frac{4}{3}\pi \rho_0 (\lambda_J/2)^3$, which is the mass contained within a sphere with a diameter equal to the Jeans length. Notice that the Jeans length is, up to a constant of order unity, equal to $c_s t_{\text{dyn}}$. i.e., it is the length over which a sound wave can travel in a dynamical time. At scales larger than the Jeans length, a sound wave cannot communicate information rapidly enough to adjust to a dynamical perturbation.

We could have derived the Jeans length by a back-of-the-envelope calculation as follows: the requirement of stability is that of hydrostatic equilibrium: $dP/dr = -G\rho M/r^2$. Setting $dP/dr \sim -P/r$, $c_s^2 \sim P/\rho$, and $M \sim \rho r^3$ we find $r = \lambda_J = c_s / \sqrt{G\rho}$.

When a fluctuation is Jeans-unstable, the fluctuation will collapse. In the limit where pressure-forces are negligible ($c_s^2 \rightarrow 0$), the solution is exponential growth:

$$\delta(t) = \delta_0 e^{t/t_{\text{dyn}}} \quad (9.15)$$

where the timescale for collapse is the dynamical time. *Perturbations grow exponentially in a static, pressureless medium.* The collapse is eventually halted by pressure. In the case of baryons the temperature will eventually rise to provide sufficient support (relevant for e.g., the collapse of gas clouds); in the case of dark matter the collapse is halted by the conversion of coherent motion into random motion; i.e., virialization.

The above discussion is quite general and relevant for the collapse of material within galaxies, for example. We will return to the Jeans length when we discuss star formation within galaxies. However, in the present discussion, we have left out one critical factor: the expansion of the universe. This one additional ingredient fundamentally changes the growth of fluctuations.

Consideration of the growth of fluctuations in an expanding universe requires several modifications to the equations above. First we will work in comoving coordinates, $\mathbf{r}$, so that the physical coordinates are $\mathbf{x} = a(t)\mathbf{r}$ where $a(t)$ is the scale factor. Note that differentials in comoving units will be different than differentials in physical units. We denote the former with a subscript c , so that $\nabla_c = a(t)\nabla$, where the unscripted differential is presumed to be with respect to physical coordinates. In an expanding universe the background velocity is not zero, instead it is governed by Hubble flow: $\mathbf{v} = H\mathbf{x} + \delta\mathbf{v} = H\mathbf{x} + a(t)\mathbf{u}$, where $\delta\mathbf{v}$ is referred to as the peculiar velocity (deviation from Hubble flow) and $\mathbf{u}$ is the comoving peculiar velocity. Moreover, the background density, $\bar{\rho}$ evolves with time. Finally, we need to switch the Lagrangian derivative, d/dt , to a derivative in the comoving frame:

$$\frac{d}{dt} = \frac{\partial}{\partial t} + \mathbf{v} \cdot \nabla = \frac{\partial}{\partial t} + H\mathbf{x} \cdot \nabla + \delta\mathbf{v} \cdot \nabla = \frac{d_c}{d_c t} + \mathbf{u} \cdot \nabla_c \quad (9.16)$$

In the equations below the $\dot{x}$ notation will refer to a time derivative in the comoving frame. We now can perform exactly the same operations as before: putting perturbed quantities into the fluid and Poisson equations, subtracting the homogeneous solution, keeping only linear terms, and combining the three equations by differentiating the continuity equation and taking the divergence of Euler's equation. The algebra is a bit more tedious, so I'll save us all the pain and jump straight to the result (for details, see Longair Ch 11.2, MvdBW Ch 4.1.1 or BT Ch 9.1.2):

$$\boxed{\ddot{\delta} + 2H\dot{\delta} = \frac{c_s^2}{a^2}\nabla_c^2\delta + 4\pi G\bar{\rho}\delta} \quad (9.17)$$

Notice that this equation is similar to Eq. 9.12, with one very important exception: the term $2H\dot{\delta}$, sometimes called the “Hubble friction” term. Further, note that the background density in the above expression is a function of a (unlike in Eq. 9.12 where it is a constant).

If we now consider solutions of the form: $\delta = D(t)e^{-i\mathbf{k}\cdot\mathbf{r}}$, where $D(t)$ is called the growth function, then we have:

$$\ddot{D} + 2H\dot{D} = D(4\pi G\bar{\rho} - c_s^2 k^2/a^2) \quad (9.18)$$

In order to make further progress we now need to assume a cosmological model, which provides the functional forms for $\bar{\rho}(t)$, $a(t)$, and $H(t)$. We will consider four illustrative examples.

- Assume a flat, matter-dominated universe filled with pressureless matter. This means $\Omega = \Omega_m = 1$ and $c_s = 0$. In such a universe we have $H(t) = \frac{2}{3t}$ and $8\pi G\bar{\rho} = 3H^2 = \frac{4}{3t^2}$. Hence:

$$\ddot{D} + \frac{4}{3t}\dot{D} - \frac{2}{3t^2}D = 0 \quad (9.19)$$

Try a solution of the form: $D \propto t^n$. This gives $n(n-1) + 4n/3 - 2/3 = 0$, which has solutions $n = 2/3$ and $n = -1$. Recall that a matter-dominated universe has $a \propto t^{2/3}$, so the growing mode solution has $\boxed{D \propto a}$.

The importance of this result cannot be over-stated. Fluctuations grow as a power-law, not exponentially. This means that the initial conditions of our universe are *not* washed away during the evolution of the universe, and it is this fact that allows us to learn about the spectrum of initial fluctuations by observing the present-day large scale structure.

This result also presents a problem: the CMB shows fluctuations at $z \sim 10^3$ of $\delta \sim 10^{-5}$. Linear growth would develop those fluctuations to only $\delta \sim 10^{-2}$ today – and yet we observed highly non-linear structures all around us. To get collapsed structures today would require fluctuations in the CMB of order 10^{-3} , two orders of magnitude larger than observed (remember: once we get to $\delta \sim 1$, collapse will proceed rapidly). The

issue is that the story we have been telling is of a baryon-only universe. As we will see in the next chapter, the addition of dark matter resolves this tension.

- Let's now consider an empty universe in which $\Omega_{\text{tot}} = 0$ and $\Omega_m = 0$. In such a universe we have $H = t^{-1}$, and

$$\ddot{D} + \frac{2}{t}\dot{D} = 0 \quad (9.20)$$

which admits solutions of the form $D = t^{-1}$ and $D = \text{constant}$. In the future, when Λ dominates and $\Omega_m \rightarrow 0$, modes will stop growing. This makes sense - the driving force of the perturbations (the matter) will be gone in the future.

- In a radiation-dominated flat universe we have $H = \frac{1}{2t}$, $a \propto t^{1/2}$, and $32\pi G\bar{\rho}/3 = t^{-2}$. In this case we have solutions of the form $D \propto a^2 \propto t$. The growth is faster than in a matter-dominated universe because the Hubble friction term is smaller.
- Finally, we can consider a universe with $\Omega_m < 1$. In this case we have:

$$\ddot{\delta} + 2H\dot{\delta} - \frac{3}{2}\Omega_m H_0^2 a^{-3}\delta = 0 \quad (9.21)$$

in which the growing mode solution is of the form:

$$f(\Omega_m) \equiv \frac{1}{H} \frac{\dot{D}}{D} = \frac{d \ln D}{d \ln a} \approx \Omega_m^{0.6} \quad (9.22)$$

i.e., modes grow slower when $\Omega_m < 1$.

A key takeaway is that fluctuations grow slowly (as a power-law) while they are small. As we will see in later chapters, once $\delta > 1$ the collapse to a fully non-linear structure is rapid, and objects with $\delta \gg 1$ are therefore not particularly useful for constraining the initial conditions of the universe. However, much of the universe today, and a greater fraction of the universe at earlier epochs is in the regime $\delta < 1$, and those structures are therefore valuable probes of the early universe. The fluctuations in the CMB are small at all scales, and so the CMB is an enormously powerful probe of the initial conditions of density fluctuations in our universe.

In this chapter we have focused on the *evolution* of an initial density fluctuation, but we have said nothing regarding the *origins* of these fluctuations. This is still an outstanding question in cosmology, and is perhaps related in some way to quantum gravity and the subsequent inflationary epoch.

Supplementary Material:

We have so far considered evolution of the density field, but what about the velocities? Recall that $\mathbf{u}$ is the comoving peculiar velocity and $\delta\mathbf{v} = a\mathbf{u}$. If we consider a velocity perturbation with only a curl component, i.e., $\nabla \cdot \mathbf{u} = 0$, then by the continuity equation densities will not grow with time, i.e., $\dot{\delta} = 0$. Moreover, Euler's equation in an expanding universe becomes: $d\mathbf{u}/dt + 2H\mathbf{u} = 0$ (our assumption of a divergence-free velocity field means that the velocity is perpendicular to the potential and so the only direction in which the velocity is non-zero is where $\nabla\Phi = 0$). Substituting $H = \dot{a}/a$ leads to $d\mathbf{u}/\mathbf{u} = -2da/a$, i.e., $u \propto a^{-2}$, or $\delta\mathbf{v} \propto a^{-1}$. The conclusion is that primordial vorticities (a curl component) decay as the universe expands. We expect the velocity field at late times to be curl-free (at least in the linear regime). This result can be seen as a statement of angular momentum conservation in an expanding universe ($mvr = \text{constant}$).

In the linear approximation, the continuity equation reads:

$$\nabla_c \cdot \mathbf{v} = -a\dot{\delta} = -a\delta \frac{\dot{D}}{D} = -a\delta H f(\Omega_m) \quad (9.23)$$

where recall that $f = d \ln D / d \ln a \approx \Omega_m^{0.6}$. The key point is that the *divergence of the velocity field is directly proportional to the density field*. This is important: the peculiar velocities, which in principle are observable, are directly linked to the density field. This idea, sometimes referred to as “cosmic flows” was very popular in the 1990s. Unfortunately, the method is beset by many systematic uncertainties and was ultimately abandoned as a cosmological tool (see Dekel 1994 for a review).

Chapter 10

Time and Scale-Dependence

In the previous chapter we worked out the linear growth of fluctuations in an expanding universe. We also introduced the concept of the Jeans length, which quantifies the critical scale at which pressure gradients are just able to balance gravity. We will now put these concepts to work and will discuss how structure in the universe develops as a function of time and scale. The story is relatively complicated because there are many physical processes to consider, although each individual process is straightforward to understand. We will begin with an introduction to the framework commonly used to describe the large scale structure in the universe.

So far we have been talking about individual fluctuations in real space, $\delta = \delta(\mathbf{x}, t)$. In reality, there are fluctuations on all scales with a range of amplitudes, and even at a given scale there will be many fluctuations with different phases. It is often more convenient to work in Fourier space, and to consider the Fourier transform of the real-space density field:

$$\delta_{\mathbf{k}} = \int d^3x \delta(\mathbf{x}) e^{-2\pi i \mathbf{k} \cdot \mathbf{x}} \quad (10.1)$$

where each $\delta_{\mathbf{k}} = A_{\mathbf{k}} e^{i\theta_{\mathbf{k}}}$ is a complex number with a phase, $\theta_{\mathbf{k}}$, and amplitude $A_{\mathbf{k}}$. (Be aware that there are different conventions regarding where to place the 2π factors in Fourier transforms.) Recall that the spatial scale associated with wavenumber k is $\lambda = 2\pi/k$.

If the phases $\theta_{\mathbf{k}}$ are uncorrelated then the field is a **Gaussian random field**. This is a key assumption in cosmology, although it can and has been tested in various ways. With this assumption the 1-pt probability distribution function for δ is a Gaussian:

$$P(\delta) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\delta^2/2\sigma^2} \quad (10.2)$$

where the variance is:

$$\sigma^2 \equiv \langle \delta^2 \rangle = \int_0^\infty 4\pi k^2 P(k) dk \quad (10.3)$$

and $P(k)$ is the **power spectrum**, which quantifies the rms fluctuation as a function of scale (whether or not the field is Gaussian), i.e.:

$$P(k) \equiv \langle A_{\mathbf{k}} A_{\mathbf{k}}^* \rangle \quad (10.4)$$

The power spectrum is a very widely used way to quantify the density field in cosmology. In linear theory each mode evolves independently, so that $\delta_{\mathbf{k}}(t) \propto D(t)$ (we know this because the growth rate, $D(t)$, derived in the previous chapter does not depend on k). This implies that the shape of the power spectrum is preserved and the amplitude grows as $D^2(t)$ (modulo the effects described below). As the field evolves beyond linear-order the probability distribution will become non-Gaussian.

We are often interested in the density field averaged (or smoothed) over some scale, R . In practice one must adopt a *window function* for smoothing, the most common choices being a top-hat or a Gaussian, where the latter is $W(r) = e^{-r^2/2R^2}$. Remember that the density field will have positive and negative fluctuations (at least until non-linear features develop), so an average over some scale will be zero:

$$\langle \Delta_R \rangle = \frac{3}{4\pi R^3} \int d^3r W(r) \langle \delta(r) \rangle = 0 \quad (10.5)$$

The rms will however not be zero: $\sigma(R) = \langle \Delta_R^2 \rangle^{1/2}$, or, in Fourier space:

$$\sigma^2(R) = \int_0^\infty 4\pi k^2 P(k) W_k^2(kR) dk \quad (10.6)$$

where $W_k(kR)$ is the Fourier transform of the window function. We can use $\sigma(R)$ measured at a particular scale to define (or measure) the normalization of the power spectrum. This is often done, and the scale chosen is 8 Mpc, with notation σ_8 . This scale is chosen in part because the rms fluctuation measured within spheres of 8 Mpc radius turns out to be close to one. There is also an observational constraint — one requires many independent spheres

in a given survey volume to accurately measure σ_8 ; this constraint limits one (or at least used to limit one) to relatively small sphere sizes.

The **2-pt correlation function** is the Fourier transform pair of the power-spectrum:

$$\xi(\mathbf{x}) = \int d^3k P(k) e^{-2\pi i \mathbf{k} \cdot \mathbf{x}} \quad (10.7)$$

The correlation function has an intuitive explanation: it quantifies the excess probability (above random) that an object (e.g., a galaxy) is found at a distance r from another object:

$$\xi(r) = \left\langle \frac{\rho(r)\rho(0)}{\rho_b^2} \right\rangle - 1 = \langle \delta(r)\delta(0) \rangle \quad (10.8)$$

where the argument is a scalar (instead of a vector) because the universe is statistically isotropic.

We will return to the statistical properties of galaxies in a later chapter.

Note that for a Gaussian field there is no additional information in the higher-order statistics. Conversely, one can compare the measured and predicted higher-order statistics to confirm that the field is Gaussian. In the highly evolved density field of the present epoch, higher-order moments do provide additional information.

It is common to adopt a power-law form for the power spectrum: $P(k) \propto k^n$, where $n = 1$ is known as the Harrison-Zeldovich spectrum, also known as the scale-invariant, or scale-free spectrum. Deviation from $n = 1$ is known as a “tilt”, and “running” of the spectral index n refers to changes in n with scale, i.e., dn/dk . Various models of inflation predict tilts and running of the spectral index, so a lot of effort goes into measuring n very precisely. The current state-of-the-art measurement is $n = 0.9665 \pm 0.0038$, as measured by the Planck collaboration in 2018. Notice the deviation from $n = 1.0$ is significant at the $\approx 8\sigma$ level.

We are now going to spend some time describing the growth of δ as a function of time and scale. There is nothing conceptually profound here, but there are a lot of moving pieces, so the details can get complicated. In particular, there are at least four important aspects to the problem: 1) there are three relevant fluid components (dark matter, baryons, and radiation); 2) components are coupled in various ways (e.g., the baryons and radiation are collisionally coupled so long as the baryons are ionized); 3) there is dissipation; 4) modes of different scales enter the horizon at different times.

We’re going to start out by neglecting the role of dark matter.

Consider the Jeans mass: $M_J \sim \rho_m \lambda_J^3$ where $\lambda_J = \sqrt{\frac{\pi c_s^2}{G\rho}}$ is the Jeans length and c_s is the sound speed. The Jeans mass is the amount of matter contained within a sphere of diameter equal to the Jeans length.

We have worked out the behavior of the densities (both matter and radiation) with redshift in previous chapters. We have not discussed in detail the sound speed, so now we must do so. For an adiabatic fluid, the sound speed is defined as:

$$c_s^2 = \left. \frac{\partial P}{\partial \rho} \right|_S \quad (10.9)$$

where the derivative is taken at constant entropy. When the universe is radiation dominated we have the usual relation for relativistic fluids: $c_s^2 = c^2/3$. When the universe is matter-dominated the sound speed is dramatically smaller and set by the temperature of the baryons ($c_s^2 \propto T$). During the transition between these two epochs the pressure is provided entirely by the radiation and the temperatures of both fluids are the same, but the inertia is provided by both fluids. We can therefore write the sound speed as:

$$c_s^2 = \frac{(\partial P / \partial T)_r}{(\partial \rho / \partial T)_r + (\partial \rho / \partial T)_m} = \frac{c^2}{3} \frac{4\rho_r}{4\rho_r + 3\rho_m} \quad (10.10)$$

We are now going to compute the evolution of the Jeans mass through the radiation-dominated and matter-dominated epochs. See Fig. 10.1 for a sketch of the key variables.

In the radiation-dominated era we have $c_s^2 = c^2/3$. The density in the Jeans length expression is given by the *total* density, which at this epoch scales as $\rho_r \propto (1+z)^4$. Meanwhile, the density in the Jeans mass expression is the baryonic (matter) density, which scales as $(1+z)^3$. The discord comes about because we are interested in asking which baryonic structures are forming, but the Jeans length is a balance between pressure forces and gravity, and the latter is sensitive to the *total* density. The upshot is that the Jeans mass goes as $M_J \propto (1+z)^{-3} \propto a^3$ in the radiation-dominated era. Numerically, we have:

$$M_J \approx 10^{29} (1+z)^{-3} \Omega_b h^2 M_\odot \quad (10.11)$$

Further, notice that $\lambda_J \sim c\sqrt{1/G\rho}$ and the horizon length $r_H \propto ct \sim c\sqrt{1/G\rho}$, i.e., the Jeans length during the radiation-dominated era is comparable to the horizon scale. This result is important, as it tells us that *modes within the horizon do not grow during the radiation-dominated era*. Instead, a perturbation is stable against gravitational collapse and

becomes a sound wave oscillating at constant amplitude. This result is due to the sound speed in the radiation fluid being comparable to the speed of light.

After matter-radiation equality, but before recombination, we have a situation in which the radiation is providing the pressure but the baryons are providing the inertia, and so we can approximate the sound speed by:

$$c_s = c \left(\frac{4\rho_r}{9\rho_m} \right)^{1/2} \approx 10^8 \left(\frac{1+z}{\Omega_b h^2} \right)^{1/2} \text{ cm s}^{-1} \quad (10.12)$$

Plugging this into our expression for the Jeans mass leads to: $M_J = 4 \times 10^{15} (\Omega_b h^2)^{-2} M_\odot$, i.e., a *constant* Jeans mass from matter-radiation equality to recombination.

Just after recombination, the pressure is provided by the thermal content of the baryons, which is *much* less than the radiation. At recombination we have $T \approx 3000\text{K}$, leading to a Jeans mass of $M_J \approx 10^5 (\Omega_b h^2)^{-1/2} M_\odot$.

In the post-recombination, matter-dominated universe, baryons will cool adiabatically (until reionization), with adiabatic index $\gamma = 5/3$, meaning $T \propto \rho^{2/3}$. Hence the Jeans mass will scale as $M_J \propto \rho^{1/2} \propto (1+z)^{3/2}$.

The next important physical process to consider is **dissipation**, which reflects the fact that the coupling between matter and radiation is not perfect. The mean free path (mfp), λ_{mfp} , is set by the Thomson scattering cross-section, σ_e : $\lambda_{\text{mfp}} = (n_e \sigma_e)^{-1}$ (the electrons are coupled to the protons via the electrostatic force). The basic idea is that any perturbation on a scale that is much smaller than the mfp will be washed away: the radiation is providing the restoring force, but it must be able to communicate to the baryons on the scale of the perturbation. In detail, the scale of interest is set by the diffusion length, $r_D = \sqrt{Dt}$, where t is cosmic time and D is the diffusion coefficient: $D = \frac{1}{3} \lambda_{\text{mfp}} c$. We can then compute a damping mass:

$$M_D = \frac{4}{3} \pi r_D^3 \rho_b = 10^{26} (1+z)^{-9/2} M_\odot \quad (10.13)$$

where to obtain the above expression we have used the radiation-dominated cosmic time-redshift relation and the usual scaling of the electron density and background density with redshift. This process is known as **Silk damping** and the associated mass is known as the Silk mass, after Joe Silk who first explored this effect (Silk 1968). By $z = 10^3$ the Silk mass is $M_D \approx 10^{12} M_\odot$ — scales smaller than this (i.e., the scale of galaxies) were washed away. In this picture, smaller scales must form via fragmentation of larger scales. This is

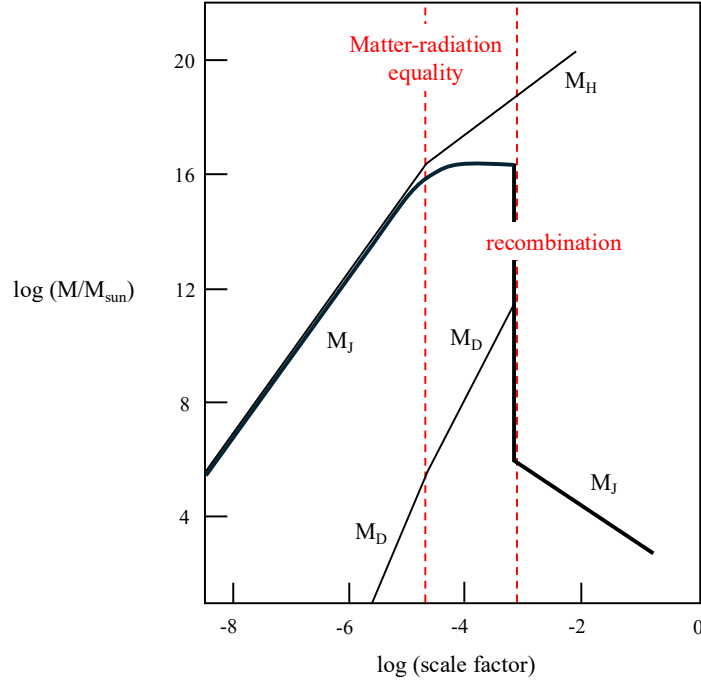


Figure 10.1: Behavior of the Jeans mass, M_J , the horizon mass, M_H , and the Silk mass M_D , vs scale factor. Adapted from Kolb & Turner (their Fig 9.3) and Longair. Note that the cosmology used to compute these quantities is not the concordance Λ CDM model, and so matter-radiation equality occurs at a much earlier epoch.

a “top-down” model of structure formation (although recall that we have ignored the dark matter so far, and inclusion of that component will change this story in an important way – see below).

Finally, we can compare these scales to the horizon scale, $M_H = \frac{4}{3}\pi r_H^3 \rho_b$ where $r_H \propto ct$. When the universe is radiation dominated we have $r_H \propto a^2$ while when the universe is matter-dominated we have $r_H \propto a^{3/2}$. This implies that $M_H \propto a^3$ in the radiation-dominated era and $M_H \propto a^{3/2}$ in the matter-dominated era.

Fig. 10.1 provides a sketch of the behavior of M_H , M_D , and M_J as a function of scale factor. The key points to remember are: 1) modes at scales below M_D are erased; 2) modes within the horizon and where $M_D < M < M_J$ oscillate; 3) modes within the horizon and with $M > M_J$ collapse. Critically, this picture is baryon-only, which the universe is not. We have spent so much time on this toy model because it illuminates the key physical concepts that we will appeal to in our discussion of structure formation including dark matter (this baryon-only model also has significant historical value).

There are two additional interesting regimes to consider. Super-horizon modes are those

where the wavelength of the perturbation is $\lambda > r_H$. It turns out that these modes do grow in the expected way dictated by linear theory, even though they are not in causal contact. This result is not obvious, but it can be sketched out by employing the Friedmann equation and treating the super-horizon fluctuations as island universes. One ends up with the island universe perturbations scaling as $\delta \propto a^2$ in the radiation-dominated era and $\delta \propto a$ in the matter-dominated era, exactly as predicted by linear theory.

Finally, note that since $t_{\text{dyn}} \propto \rho_m^{-1/2}$ while $t_{\text{univ}} \propto \rho_{\text{tot}}^{-1/2}$, when $\rho_r \gg \rho_m$ we have $t_{\text{univ}} \ll t_{\text{dyn}}$. This implies that modes can't grow at very early epochs even if they are Jeans-unstable. The universe is expanding too rapidly for the mode to collapse. This is known as the Meszaros effect.

We are now going to consider how this picture changes with the addition of dark matter. Because dark matter is not strongly coupled to the other fluids (except via gravity), we can expect the dynamics to be rather different. See Fig. 10.2.

A key concept in the dark matter fluid is the **free streaming length**, which is the distance a particle will travel in a Hubble time (set by the particle velocity). Qualitatively, the idea is that the particles can travel in random directions up to a length λ_{FS} and any perturbation on scales $\lambda < \lambda_{\text{FS}}$ will therefore be erased. Although the physics is somewhat different, the end result is similar to Silk damping. The relevant velocity for the dark matter particles is set during freeze-out (an event similar to recombination but for particle creation). For very light particles, such as neutrinos, the velocity at freeze-out is relativistic, and so all small scale structure is washed out. (For neutrinos to supply the dark matter they would need a rest mass of $10 - 30$ eV; the actual masses are $\lesssim 1$ eV, so they cannot.) This class of model is generically called hot dark matter (HDM), and is firmly ruled out by observations of the power spectrum at small scales. In contrast, cold dark matter (CDM) refers to massive particles that were non-relativistic at freeze-out. This category of particles is called WIMPs, for weakly-interacting massive particles (rest mass $10 - 10^3$ GeV), and their free-streaming length is very small, so all cosmologically-interesting scales are preserved. There is also an intermediate class of models called warm dark matter (WDM; rest mass ~ 1 keV), which has some attractive properties and is not ruled out by observations. The “cold”, “warm”, and “hot” simply refer to the velocities of the dark matter particles at freeze-out. See Chapter 13 for details.

After matter-radiation equality, but before recombination, dark matter perturbations within the horizon can grow, while baryonic perturbations cannot, because the latter are coupled

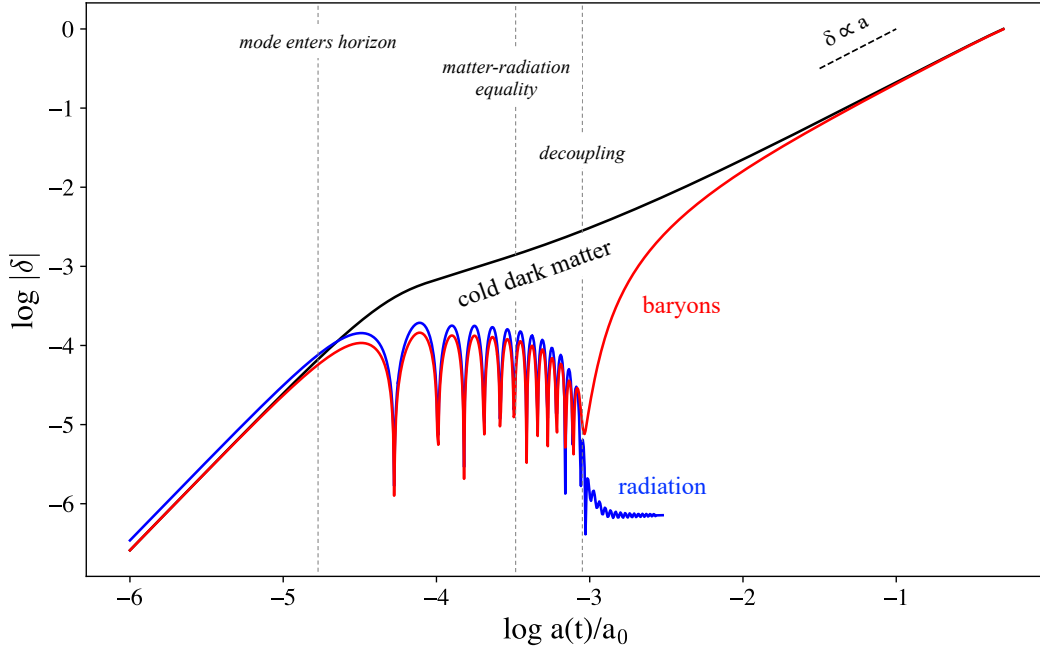


Figure 10.2: Growth of density fluctuations with scale factor for a Λ CDM cosmology and $k = 0.3 \text{ Mpc}^{-1}$.

to the photons, which are providing a restoring force. This means that the dark matter perturbations will have larger amplitudes by the time of recombination. Once the baryons decouple, their perturbation amplitudes will quickly “catch up” to that of the dark matter. This means that the observed fluctuations in the CMB are *not* a direct reflection of the underlying matter fluctuations.

Finally, as in the baryon-only picture, a dark matter perturbation that enters the horizon before matter-radiation equality will not grow. The reason is simple: before matter-radiation equality the energy density is dominated by the radiation, and so the timescale of the Universe is much shorter than the dynamical time of a perturbation.

Let’s return to the power spectrum, $P(k)$. We asserted earlier that the initial power spectrum is a power-law of the form $P(k) \propto k^n$ with $n \approx 1$. However, the *observed* power-spectrum is strongly affected by the effects we have discussed in this chapter. One therefore often writes the power spectrum as:

$$P(k, t) = P(k, t = 0) T^2(k) D^2(t) \quad (10.14)$$

where $P(k, 0)$ is the initial power spectrum, $D(t)$ is the linear growth factor, and $T(k)$ is the transfer function. This last term describes the growth of fluctuations after they re-enter the horizon (we say “re-enter” because these modes were presumably within the horizon before inflation). Its functional form is quite complicated and modern analysis uses numerical tools to evaluate the transfer function (e.g., [CAMB](#)).

We can describe one important feature in the transfer function. For modes that enter the horizon before matter-radiation equality, their amplitudes are stalled relative to those modes which have not yet entered the horizon. The evolved power spectrum will therefore be suppressed on small scales, with the deviation from $P(k) \propto k^n$ occurring up to the largest scale where the mode has just entered the horizon at equality: $\lambda \sim ct_{\text{eq}}/a_{\text{eq}} \approx 35$ Mpc (in comoving units).

Modes enter the horizon when $kct/a \approx 1$. In the radiation-dominated era $t \propto a^2$, so a given mode k will enter the horizon when $a_{hc} \propto k^{-1}$. During the radiation-dominated era the growth factor goes as $D \propto a^2$; modes beyond the horizon continue to grow at this rate, while the ones within the horizon stall. So the modes that have entered the horizon are damped as k^{-2} , and the power spectrum is therefore suppressed by k^4 on small scales. This means that on small scales we can expect $P(k) \propto k^{n-4} \propto k^{-3}$.

The importance of Eq. 10.14 is that we can infer the initial conditions based on measurement of $P(k, t)$ at any scale and epoch. The only requirement is that the observed scale is still in the linear regime at the epoch of the measurement.

Chapter 11

Non-Linear Evolution & Spherical Collapse

In the previous two chapters we worked out the linear growth of fluctuations and their consequences. There is a large literature developing non-linear perturbation theory in order to describe the growth of fluctuations beyond linear order. However, once a perturbation goes non-linear (i.e., once it can no longer be described by linear theory), the evolution is quite rapid. As a consequence, non-linear theory is of limited utility. In this chapter we will discuss the non-linear evolution of fluctuations and will focus on the spherical collapse model for the growth of fully non-linear structures. By the end of this chapter we will have a good idea for the scale and density of collapsed objects, and the timescales involved.

Let's begin by considering which scales become non-linear at which epochs. Let's assume a matter-dominated universe, in which $\delta \propto a \propto t^{2/3}$ and $\rho_b \propto t^{-2}$, where ρ_b is the background matter density (the universe was matter-dominated during most of the evolution of the universe, from $z \sim 10^3$ to $z \sim 0.5$, so this is a good assumption).

Recall from the last chapter that the initial power spectrum is assumed to be a power-law: $P(k) \propto k^n$, and that the variance of the fluctuations averaged on a physical length scale R at a fixed time is $\sigma^2(R) \propto \int k^2 dk P(k) \propto k^{n+3} \propto R^{-(n+3)}$. Because the variance of fluctuations is defined as $\sigma_k^2 \equiv \langle \delta_k^2 \rangle$, we also have, given our assumption of a matter-dominated universe: $\sigma_k^2 \propto t^{4/3}$, which is the time-dependence of the fluctuation at a given scale. We can combine these two dependencies and arrive at:

$$\sigma^2(R, t) \propto t^{4/3} R^{-(n+3)} \quad (11.1)$$

We will now define the non-linear scale, R_{NL} by $\sigma^2(R_{\text{NL}}) = 1$. Given this definition, we can conclude that the time for a given scale to go non-linear scales as: $t_R \propto R^{3(n+3)/4}$.

From these simple arguments we can deduce several important consequences. Small scales will go non-linear before large scales, so long as the initial power spectrum has $n > -3$. This is sometimes referred to as “bottom-up” structure formation.

Objects collapsing at $t = t_R$ will have a density of order $\bar{\rho}(t_R)$, hence the collapsing object will have $\rho \propto t_{\text{NL}}^{-2} \propto R^{-3(n+3)/2}$. For $n > -3$, smaller scales collapse earlier, when the characteristic density of the universe was higher. Therefore, smaller objects will generally both collapse earlier and have higher characteristic densities.

This simple analysis has real world value: as we will see in the next chapter, dark matter halos, including their density profiles, can be characterized by two parameters: mass (or scale) and collapse time.

We are now going to embark on a detailed analysis of a concept called **spherical collapse**. The basic setup is straightforward: start with a small initial density perturbation, and use gravitational dynamics and the virial theorem to evolve the density through the collapse phase, into a fully bound, self-gravitating structure. While idealized, this toy model provides solid intuition for how real density perturbations evolve into the bound objects we call dark matter halos. See e.g., Gunn & Gott (1972) for an early use of the spherical collapse model to understand the properties and evolution of massive clusters of galaxies. The classic reference on this topic is Bertschinger (1985).

Suppose that there is an initial density perturbation of the form: $\rho(r, t_i) = \rho_b(t_i)[1 + \delta_i(r)]$ where ρ_b is the background density and δ_i is the initial density contrast, which is a decreasing function of r (in a coordinate system centered on the peak density).

A key assumption is that shells of material do not cross (which is a good assumption at early times, and more dubious at later times), in which case the mass contained interior to a given a shell is constant:

$$M(r, t) = M(r_i, t_i) = \rho_b(t_i) \left(\frac{4}{3} \pi r_i^3 \right) (1 + \delta_i) \quad (11.2)$$

The equation of motion for the shell is:

$$\ddot{r} = -\frac{GM(r, t)}{r^2} \quad (11.3)$$

Integrating once, we have:

$$\frac{1}{2}\dot{r}^2 - \frac{GM(r, t)}{r} = E \quad (11.4)$$

where E is an integration constant and is the total energy per unit mass of the shell. If $E < 0$ the shell eventually collapses, while if $E > 0$ the shell never collapses.

At t_i we have $\delta_i \ll 1$ so the initial overdensity is expanding with the Hubble flow with zero peculiar velocity ($v_i = 0$). Then we have, for the initial velocity of the shell: $\dot{r}_i = H_i r_i$, and the kinetic energy per unit mass of the shell is $K_i \equiv \dot{r}_i^2/2 = H_i^2 r_i^2/2$. Likewise, the potential energy per unit mass is:

$$|U_i| = \frac{GM}{r_i} = \frac{4}{3}\pi G \rho_b(t_i) r_i^2 (1 + \delta_i) \quad (11.5)$$

Recall that the background density can be written as $\rho_b = \Omega_m \rho_c$, where ρ_c is the critical density, which can be written as: $\rho_c = \frac{3H^2}{8\pi G}$. Combining these two expressions, we have:

$$|U_i| = \frac{1}{2}H_i^2 r_i^2 \Omega_i (1 + \delta_i) = K_i \Omega_i (1 + \delta_i) \quad (11.6)$$

where above and for the rest of this chapter we set $\Omega = \Omega_m$.

The total energy per unit mass of the shell is then:

$$E = K + U = K_i - K_i \Omega_i (1 + \delta_i) \quad (11.7)$$

$$= K_i \Omega_i (\Omega_i^{-1} - 1 - \delta_i) \quad (11.8)$$

which means that collapse will occur if $(1 + \delta_i) > \Omega_i^{-1}$. In a closed or flat universe we have $\Omega_i^{-1} \leq 1$, so the condition for collapse is satisfied for all positive density fluctuations. However, in an open universe δ_i must be above some critical value (greater than zero) for collapse to occur.

Now we are going to consider the dynamics of a collapsing shell ($E < 0$). Such a shell is initially expanding with the Hubble flow, and will expand to a maximum radius $r_{\max}$, at which point $\dot{r}(r_{\max}) = 0$. This point is often called turnaround. At this point the total energy of the shell is:

$$E = U_{\max} = -\frac{GM}{r_{\max}} = -\frac{r_i}{r_{\max}} |U_i| = -\frac{r_i}{r_{\max}} K_i \Omega_i (1 + \delta_i) \quad (11.9)$$

Equating this with our expression for the initial energy leads to:

$$\frac{r_{\max}}{r_i} = \frac{1 + \delta_i}{\delta_i - (\Omega_i^{-1} - 1)} \quad (11.10)$$

which means that smaller initial fluctuations have larger $r_{\max}/r_i$, i.e., they expand further before turnaround and take longer to collapse. This conforms with our intuition.

Let's return to our equation of motion: $\frac{1}{2}\dot{r}^2 - \frac{GM(r,t)}{r} = E$. For $E < 0$ the solution is a parametric equation of the form:

$$r = A(1 - \cos \theta) \quad (11.11)$$

$$t = B(\theta - \sin \theta) \quad (11.12)$$

where $A^3 = GMB^2$.

Notice that turnaround occurs at $\theta = \pi$, where $r = r_{\max} = 2A$. From Eq. 11.10 we then have:

$$A = \frac{r_i}{2} \frac{1 + \delta_i}{\delta_i - (\Omega_i^{-1} - 1)} \quad (11.13)$$

we can also work out the expression for B from its relation to A and M , along with Eq. 11.2 and the definition of the background density in terms of Ω and H :

$$B = \frac{1}{2H_i\Omega_i^{1/2}} \frac{1 + \delta_i}{[\delta_i - (\Omega_i^{-1} - 1)]^{3/2}} \quad (11.14)$$

We can now write the parametric solution (Eq. 11.11) in a more revealing form:

$$r = \frac{r_{\max}}{2} (1 - \cos \theta) \quad (11.15)$$

$$t = \frac{t_{\max}}{\pi} (\theta - \sin \theta) \quad (11.16)$$

where $t_{\max}$ is the time at which turnaround occurs.

Our next major task is to compute the evolution in density of the perturbation. This turns out to be a bit involved, but it is very illuminating. See Fig. 11.1 for a sketch of the evolution.

We will consider the mean density interior to the shell:

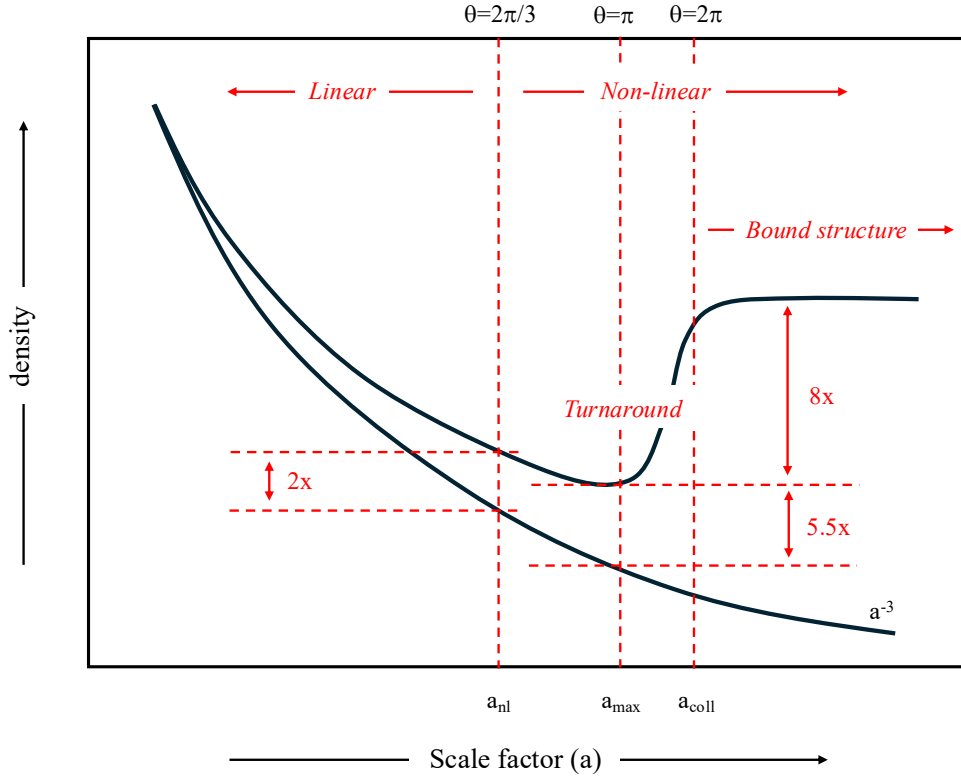


Figure 11.1: Sketch of the density evolution of a fluctuation compared to the background density (ρ_b), as a function of scale factor a . Adapted from Longair.

$$\rho(t) = \frac{3M}{4\pi r^3} = \frac{3M}{4\pi A^3 (1 - \cos\theta)^3} \quad (11.17)$$

in a flat universe the background density is: $\rho_b(t) = (6\pi G t^2)^{-1}$; we can use this to compute the evolution of the density contrast with time:

$$\frac{\rho(t)}{\rho_b(t)} = 1 + \delta(t) = \frac{3M}{4\pi A^3} (6\pi G B^2) \frac{(\theta - \sin\theta)^2}{(1 - \cos\theta)^3} \quad (11.18)$$

where we replaced t in the expression for ρ_b with the parametric solution. Using $A^3 = GMB^2$ we arrive at our final expression:

$$\delta = \frac{9}{2} \frac{(\theta - \sin\theta)^2}{(1 - \cos\theta)^3} - 1 \quad (11.19)$$

At early times, when t , θ and δ are small, we have $\delta = 3\theta^2/20$ and $t = B\theta^3/6$ and so:

$$\delta = \frac{3}{20} \left(\frac{6t}{B} \right)^{2/3} \quad (11.20)$$

Notice that this recovers the linear theory prediction that $\delta \propto t^{2/3}$. Moreover, notice that $\delta_{\text{NL}} = 1$ at $\theta = 2\pi/3$.

At turnaround, when $\theta = \pi$ and $t = \pi B$, we have $\delta_{\text{NL}} = \frac{9}{2} \frac{\pi^2}{8} - 1 = 4.55$, so $\rho/\rho_b = 5.55$. Linear theory would predict $\delta_{\text{lin}} = \frac{3}{20}(6\pi)^{2/3} \approx 1.06$. This means that the density field evolved according to linear theory is approaching turnaround when $\delta_{\text{lin}} \approx 1$. Obviously by this point the assumption of linear theory has long since broken down.

Finally, at $\theta = 2\pi$ we again have $r = 0$, i.e., the object has collapsed. At this point $t = 2\pi B$, i.e., $t_{\text{coll}} = 2t_{\text{turnaround}}$. The non-linear density (Eq. 11.19) is formally infinite, and $\delta_{\text{lin}} = \frac{3}{20}(12\pi)^{2/3} \approx 1.686$. This last number is important, as we will use it in the next chapter to connect the linearly-evolved density field to collapsed objects. In other words, although linear theory is *wrong* when δ is large, we can use the mapping provided above to identify turnaround ($\delta_{\text{lin}} = 1.06$) and collapse ($\delta_{\text{lin}} = 1.686$) purely from the linearly-evolved density field. This is useful, because it is easy to linearly-evolve the density field.

Of course, the collapse never proceeds to $\theta = 2\pi$. In practice, the assumptions of no shell crossing and negligible random velocities eventually break down. Violent relaxation will occur, enabling the transfer of coherent infall motion into random motion. Eventually the random motion will provide sufficient thermal energy to halt collapse, and the system will reach dynamical (virial) equilibrium.

Note that δ is usually a decreasing function of r , so smaller scales of a perturbation collapse first. This leads to the inside-out growth of perturbations.

Our final task is to compute the density contrast of the collapsed object.

At turnaround all of the energy is in potential energy: $E = U = -\frac{GM}{r_{\text{max}}}$, while at virialization we have (via the virial theorem): $E = U/2 = -\frac{GM}{2r_{\text{vir}}}$, which implies $r_{\text{vir}} = r_{\text{max}}/2$, i.e., the object only collapses in size by a factor of two from turnaround to virialization ($t_{\text{vir}} = 2t_{\text{max}}$). This means $\rho_{\text{vir}} = 8\rho_{\text{max}}$. To compute the density *contrast* of the collapsed object, we need to know ρ_b at virialization. We start by noting that $\rho_{\text{max}} = \frac{9}{16}\pi^2\rho_b(t_{\text{max}})$. We also know that the background density evolves as $\rho_b \propto t^{-2}$, so we have

$$\rho_b(t_{\max}) = \rho_b(t_{\text{vir}}) \left(\frac{t_{\text{vir}}}{t_{\max}} \right)^2 = 4\rho_b(t_{\text{vir}}) \quad (11.21)$$

Therefore:

$$\rho_{\text{vir}} = 8\rho_{\max} = 8 \cdot 4 \cdot \frac{9}{16} \pi^2 \rho_b(t_{\text{vir}}) \quad (11.22)$$

$$\boxed{\frac{\rho_{\text{vir}}}{\rho_b} \equiv \Delta = 18\pi^2 \approx 180} \quad (11.23)$$

where Δ is the density contrast of a collapsed object. We can therefore expect a collapsed object in the universe to have a density contrast relative to the background matter density of ~ 200 . Note that Δ depends on cosmology.

The intuition from spherical collapse is used to *define* virialized regions in dark matter cosmological N -body simulations. There is an entire cottage industry of scientists attempting to define collapsed objects (what are called **dark matter halos**) in ever more sophisticated ways. Remarkably, the relatively simple spherical collapse model correctly describes the basic picture, including the quantitative details (at some level), which is why we've spent an entire chapter on it. Extensions to the above approach consider non-spherical collapse (i.e., ellipsoidal collapse).

If we *define* a virialized object as a region within which the mean density is Δ times the background density, then we can define **virial parameters** as follows (below assuming $\Omega_m = 1$):

$$R_{\text{vir}} = \left(\frac{2GM}{\Delta H^2} \right)^{1/3} = 0.8 \text{ kpc} \left(\frac{M}{10^8 M_\odot} \right)^{1/3} \left(\frac{10}{1+z} \right) \quad (11.24)$$

$$V_{\text{vir}} = \left(\frac{GM}{R_{\text{vir}}} \right)^{1/2} \propto M^{1/3} = 23 \text{ km s}^{-1} \left(\frac{M}{10^8 M_\odot} \right)^{1/3} \left(\frac{10}{1+z} \right)^{-1/2} \quad (11.25)$$

$$T_{\text{vir}} = \frac{\mu m_p V_{\text{vir}}^2}{2k} \propto M^{2/3} = 2 \times 10^4 \text{ K} \left(\frac{M}{10^8 M_\odot} \right)^{2/3} \left(\frac{10}{1+z} \right)^{-1} \quad (11.26)$$

Above are the virial radius, virial velocity, and virial temperature. Notice that since the mean density at R_{vir} is constant for all halos by definition, so is the dynamical time at the virial radius constant for all halos at a given epoch.

The above equations are referenced to small halos at high redshift. Scaling to the Milky Way, which has $M \approx 10^{12} M_\odot$ at $z = 0$, we find $R \approx 170 \text{ kpc}$, $V = 150 \text{ km s}^{-1}$ and $T \approx 10^6$

K. Note that the virial radius depends on the adopted overdensity convention; the more commonly quoted R_{200} for the Milky Way is closer to 200 kpc.

Chapter 12

Dark Matter Halos

In this chapter we turn our attention to dark matter halos. These are the natural “units” in the dark matter density field; they are fully collapsed, virialized objects and are clearly identifiable in the density field. Halos grow over time both via the smooth accretion of surrounding matter and via mergers with other halos. One reason to study dark matter halos in detail is because they are the sites where galaxies form.

There are many properties we would like to know of halos: their abundance, clustering, internal properties, and growth histories, to name a few. Until the early 2000s, computers were not fast enough to simulate the dark matter density field with the volume and resolution needed to reliably answer these questions. Partly for this reason, and partly because of the intuition afforded by analytic models, much of the early work related to halos was (semi-)analytic.

Our first goal is to determine the **mass function** of collapsed objects, i.e., the number (or number density) of halos as a function of their mass.

Recall that the density contrast, $\delta(x)$, is a Gaussian random variable. Smoothing of this field by a window function, W , of width R , i.e.,

$$\delta(x, R) = \int d^3x' W(|x - x'|; R) \delta(x') \quad (12.1)$$

will result in a new field that is also a Gaussian. Further, recall that the variance on scale R is $\sigma^2(R) = \langle \delta^2(x; R) \rangle$.

The Gaussian nature of $\delta(x, R)$ means that the 1-pt probability distribution is:

$$P(\delta; R) = \frac{1}{\sqrt{2\pi\sigma^2(R)}} e^{-\delta^2/2\sigma^2(R)} \quad (12.2)$$

Our approach is going to be to use linear theory to predict the mass spectrum of collapsed objects. This sounds like a bad idea, as collapsed objects are in the deeply non-linear regime. However, our derivation of spherical collapse gave us the insight that an object has collapsed and formed a virialized object when the linearly-evolved density field reaches $\delta_{\text{coll}} = 1.69$.

For a given smoothing scale, R , we can associate a corresponding mass via $M = \frac{4}{3}\pi\rho_m R^3$ where ρ_m is the background matter density. Note that this means there is a one-to-one correspondence between smoothing scale R and mass scale M , so we can write $\sigma^2(R) = \sigma^2(M)$.

We can therefore imagine taking the density field, smoothing it by scale R (or equivalently, mass M), evolving the field forward in time according to linear theory via $\delta(x, t) = \delta(x, 0)D(t)$, and then calling that mass a collapsed object when $\delta(R) > \delta_{\text{coll}}$. This is the basic approach taken by Press & Schechter (1974), in what is now called the **Press-Schechter formalism**.

The Press-Schechter ansatz states that the fraction of mass in halos of mass $> M$ is given by the fraction of the Gaussian distribution $\delta(R)$ that exceeds $\delta_{\text{coll}} = 1.69$, i.e.,

$$F(> M) = \int_{\delta_{\text{coll}}}^{\infty} P(\delta, R) d\delta = \frac{1}{\sqrt{2\pi}} \int_{1.69/\sigma}^{\infty} d\nu e^{-\nu^2/2} \quad (12.3)$$

where we have defined $\nu \equiv \delta_{\text{coll}}/\sigma$. Notice that ν encodes (or hides) a lot of crucial information via its dependence on σ . In particular, ν encodes the combination of the primordial power spectrum and the subsequent linear theory growth of the density field. In particular, $\sigma^2 \propto D^2(t)$, so as time progresses ν , which sets the barrier for collapse, decreases.

The number density of halos is then:

$$\frac{dn}{dM} = \frac{\rho_m}{M} \frac{dF}{dM} \quad (12.4)$$

where the first term on the right hand side converts from mass to number density, and the second is the differential mass fraction. We then have:

$$\frac{dn}{dM} = \frac{\rho_m}{M} \frac{1}{\sqrt{2\pi}} \frac{d\nu}{dM} e^{-\nu^2/2} \quad (12.5)$$

or, equivalently:

$$\frac{dn}{d \ln M} = \frac{\rho_m}{M} \frac{1}{\sqrt{2\pi}} \left| \frac{d \ln \nu}{d \ln M} \right| \nu e^{-\nu^2/2} \quad (12.6)$$

we have expressed the mass function in this way for the following reason. Recall that the power spectrum is $P(k) \propto k^n$ and so $\sigma^2(R) \propto R^{-(3+n)}$. By definition $\nu \propto \sigma^{-1}$, which means $\nu \propto M^{(3+n)/6}$. Therefore:

$$\left| \frac{d \ln \nu}{d \ln M} \right| = \frac{|n+3|}{6} \quad (12.7)$$

i.e., this term encapsulates the dependence on the initial power spectrum.

Now we come to the most dubious step in the Press-Schechter formalism. We have only considered the positive half of the Gaussian, i.e., $F(0) = 1/2$, so that only half of the mass is in collapsed objects. The “solution” proposed by PS was to simply multiply our mass function by two. The final result being the famous Press-Schechter mass function:

$$\boxed{\frac{dn}{d \ln M} = \frac{\rho_m}{M} \sqrt{\frac{2}{\pi}} \frac{|n+3|}{6} \nu e^{-\nu^2/2}} \quad (12.8)$$

Eq. 12.8 is probably the most important single result in low-redshift ($z \lesssim 20$) cosmology.

The factor of two swindle has led to much hand-wringing and post-hoc justification. One can obtain a more satisfactory result by considering the fact that locally under-dense regions can reside within a larger scale overdensity and so the under-dense region has some probability of collapsing on the larger scale. This is one example of the “cloud-in-cloud” problem (see Bond et al. 1991 for a sense of the complexity involved).

The resulting mass function scales as $dn/dM \propto M^{-2}$ at low masses and has an exponential cut-off at a mass scale denoted M_* . At $z = 0$ this scale is approximately $10^{13} M_\odot$. To give some examples: galaxy clusters have total masses $M > 10^{14} M_\odot$, which corresponds to $\nu = 2$ and hence $dn/d \ln M = 10^{-5} \text{ Mpc}^{-3}$, while typical Milky Way-like galaxies have $M = 10^{12} M_\odot$, corresponding to $\nu \approx 0.7$ and a number density of $dn/d \ln M \approx 3 \times 10^{-3} \text{ Mpc}^{-3}$.

The power-law behavior at low mass (low ν) is close to -2 because the logarithmic derivative of ν and ν itself do not vary strongly with mass for small masses. This means that there are a formally divergent number of small halos: $\int \frac{dn}{dm} dm = \int m^{-2} dm \sim m^{-1}$, although note that the total *mass* is only logarithmically-divergent. This is all a bit pedantic as there is

in fact a minimum mass to the mass function set by the free-streaming length of cold dark matter; the scale is not known but is predicted to be somewhere in the Earth-mass regime.

The behavior of the mass function at low masses is in striking contrast to the observed stellar mass function of galaxies, which has a slope closer to $\alpha \approx -1.2$. This discrepancy is the origin of many “crises” in modern cosmology. We’ll return to this point in later chapters.

Mass Functions: The Press-Schechter mass function is our first example of a mass function. We will return to this topic in subsequent chapters when we discuss galaxies and stars. Be aware that mass functions are commonly expressed as either dn/dM or $dn/d\ln M$ and they will therefore differ by one power of M . The latter is often regarded as the more “natural” unit, i.e., the number per log mass interval.

Press-Schechter theory is not perfect but does get the basic correct behavior at the factor of ~ 2 level when compared to modern simulations (see Fig. 12.1). There have been a variety of improvements, mostly related to extending the framework to ellipsoidal rather than spherical collapse. The math becomes *much* more involved. See Sheth, Mo, & Tormen (2001) for an example. Modern halo mass functions are essentially simply fitting functions to large suites of simulations. This has the advantage of accuracy, but at the cost of flexibility. One of the valuable aspects of the PS mass function is that it “works” for an arbitrary initial power spectrum, so one can explore mass functions for a very wide range of cosmological models.

The value of the Press-Schechter formalism extends far beyond simply computing the halo mass function. The most influential post-PS paper was that of Bond et al. (1991), which developed the excursion set formalism (also called extended Press-Schechter, or EPS). We won’t get into the details here, but the basic idea is that one can imagine smoothing the density field on progressively larger scales (larger masses). If the smoothing is done in k -space, then the resulting function $\delta(R)$ is a *random walk* because k modes are uncorrelated in linear theory. This insight allows one to employ the considerable mathematical machinery developed for random walks in order to study the excursions of the density field as a function of smoothing scale.

EPS can show where the missing factor of two in the original formulation comes from. But the formalism can do much more. In EPS one can compute the *conditional mass function*, i.e., the mass function conditioned on a given large-scale density. EPS can also compute the clustering of halos and the growth history of halos (i.e., “merger trees”). EPS was very popular in the 1990s–2000s, when simulations were of insufficient volume and/or resolution

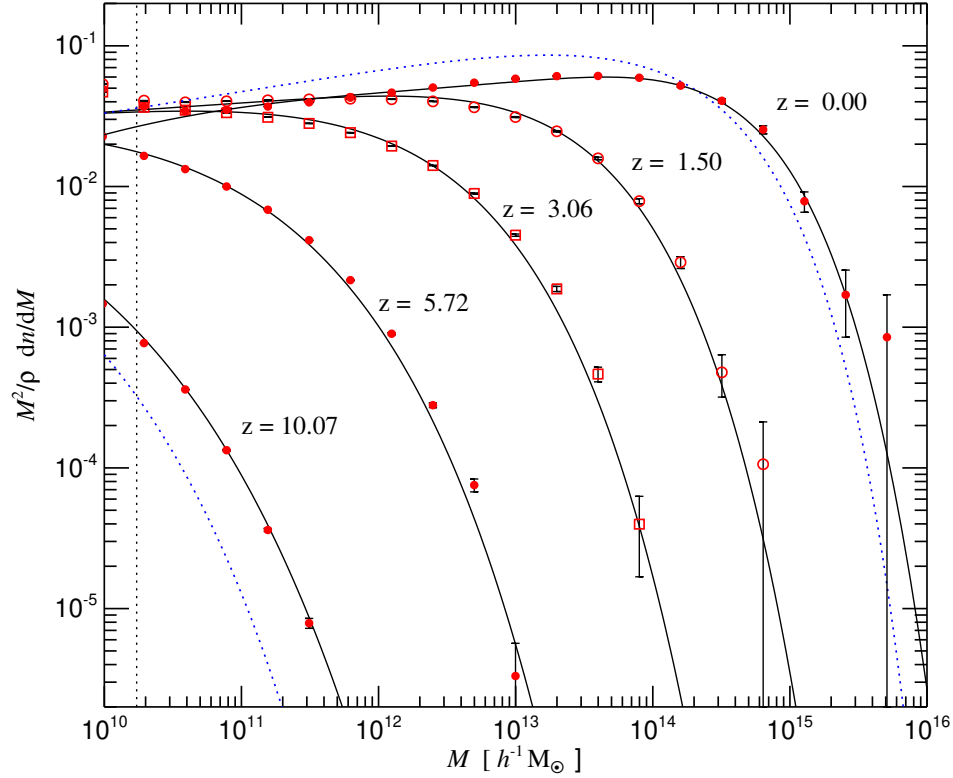


Figure 12.1: Halo mass function from the Millennium Simulation (points) compared to Press-Schechter theory (blue dotted lines) and a fitting function (solid lines). From Springel et al. (2005).

to track halos reliably. However, at present the simulations have become very reliable, so EPS has fallen somewhat out of fashion.

We will now turn to the internal structure of dark matter halos. The seminal paper on this topic is Navarro, Frenk, & White (1997), often referred to as “NFW”. They measured the density profiles of halos in simulations and found that the profiles obeyed a universal function, independent of mass and even the underlying cosmology:

$$\frac{\rho(r)}{\rho_{\text{crit}}} = \frac{\delta_c}{(r/r_s)(1 + r/r_s)^2} \quad (12.9)$$

where ρ_{crit} is the critical density from cosmology and r_s is a scale radius. Notice that the NFW density profile scales as r^{-1} at small radii and r^{-3} at large radii.

The virial radius, r_{vir} is defined such that $\bar{\rho}(r_{\text{vir}}) = \Delta \rho_{\text{crit}}$, where Δ is the threshold density often taken to be 200 (although be aware that there are many different conventions in the

literature). NFW then defined the **concentration**: $c \equiv r_{\text{vir}}/r_s$. The normalization of the density profile, δ_c is linked to the concentration via:

$$\delta_c = \frac{200}{3} \frac{c^3}{\ln(1+c) - c/(1+c)} \quad (12.10)$$

Notice that the virial mass, M_{vir} , and virial radius are directly related to one another. Therefore the mass and concentration uniquely determine the density profile of the halo, including both shape and amplitude.

From analysis of simulations we know that the concentration of halos scales as:

$$c \propto \frac{M^{-0.1}}{1+z} \quad (12.11)$$

Low-mass halos are more concentrated than high-mass halos, and concentrations at fixed mass are lower at higher redshift. At $z = 0$, $c = 10 - 20$ for MW-mass halos, while $c \sim 2 - 5$ for clusters.

Simulations indicate that the early phase of rapid accretion builds the r^{-1} profile, while the later slow accretion builds the r^{-3} profile. *When* this transition between rapid and slow accretion occurs sets the concentration. Note that the inner scale r_s changes little at late time. The concentration is growing because the virial radius is increasing.

Simulations also show that the mass accretion histories of halos are well described by $M(z) = M_0 e^{-\alpha z}$, where α is a function of concentration (Wechsler et al. 2002), such that at fixed mass, high concentration halos form earlier than low concentration halos. This conforms with our earlier discussion of spherical collapse - if a halo collapses earlier it will be denser, as it is collapsing out of a denser universe. We therefore often think of concentration as mapping directly to a formation time for the halo.

Exactly how the halo mass is assembled has been the subject of an enormous body of work. Certainly a large fraction is built directly through the accretion of smaller halos (such events are called mergers). As much as 20% of the mass growth may be due to “smooth” accretion, i.e., mass unassociated with halos. This statement is controversial, as one must always contend with the resolution limit of the simulation. Recently, attention has been drawn to *pseudo-evolution* of the halo mass as a consequence of the fact that the outer boundary (the virial radius) is defined with respect to either the critical or mean background density, and both of these evolve with redshift. So even a static, intrinsically un-evolving halo will grow in virial mass as a consequence of the definition of the virial radius.

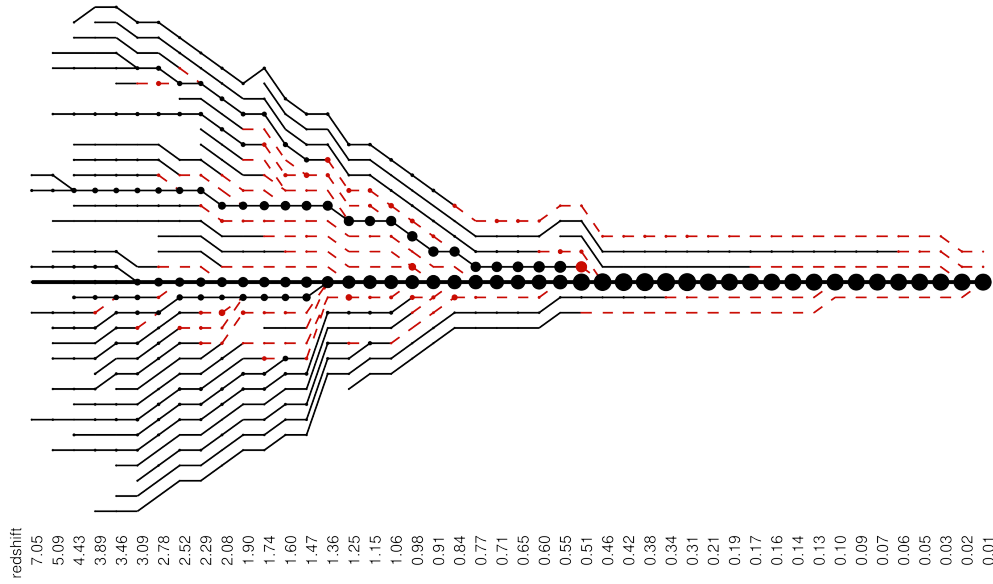


Figure 12.2: Example of a halo merger tree, with time progressing toward the right. Symbol sizes are proportional to $\log(M_{\text{halo}})$. Red dashed lines are subhalos. The resolution limit of the simulation limits the tree from extending arbitrarily far back in time. From Stewart et al. (2008).

For the discrete objects that are accreted, one can construct a **merger tree**, which quantifies which halos are accreted at which times (see Fig. 12.2). Merger trees can be constructed from EPS theory or via simulations, and they form the bedrock of semi-analytic models of galaxy formation, as we will see in later chapters and problem sets.

There is a deeper question as to the physical significance (if any) of the virial radius. The name would suggest that the material within this radius is virialized, but this is generally not the case. Material at the virial radius is often radially infalling into the halo for the first time. Recent work has attempted to isolate a physical boundary via the so-called “splash-back radius”, or the zero-velocity surface. Such a surface is defined by the turnaround radius of the most weakly bound particles. Such a surface does generally appear in simulations, at radii generally somewhat larger than the classical virial radius. However, the splash-back radius is much more difficult to define in simulations, and probably impossible in most cases in observations.

Many PhD theses have been devoted to understanding various other aspects of dark matter halos, including their shapes, angular momenta, the properties of the **subhalos** residing within larger halos (see Fig. 12.3), their velocity anisotropy, connection to the cosmic web, etc.

The subhalos are particularly important, as they are believed to be the sites of satellite

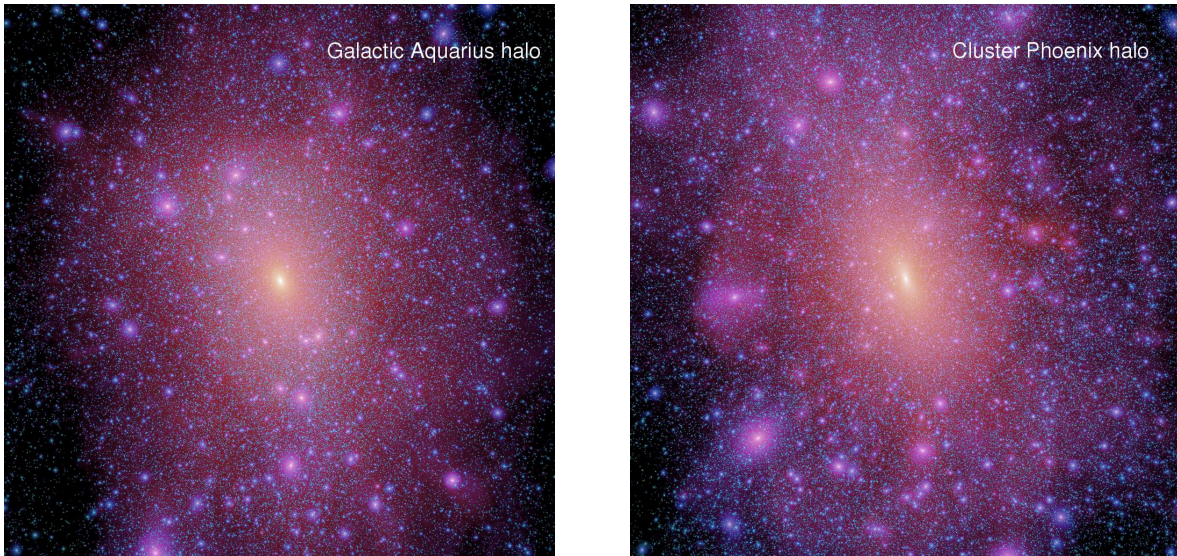


Figure 12.3: N-body simulation of CDM halos. The color is proportional to density. In both panels one sees a dominant “main” halo and hundreds of smaller subhalos that orbit within. The left panel shows a $10^{12}M_{\odot}$ halo while the right panel shows a halo of mass $6 \times 10^{14}M_{\odot}$. It is remarkable that halos differing by a factor of 600 in mass display such similar internal structure. From Springel et al. (2008); Gao et al. (2012).

galaxies orbiting within the parent system. At any given time, some 20 – 50% of the parent halo mass is composed of smaller subhalos. The subhalos have their own mass function that is nearly scale-free (independent of the parent mass; see Fig. 12.3). In detail, the properties of the subhalos, including their orbits, numbers, masses, etc. are governed by the dynamical processes we discussed in earlier chapters (tidal stripping, dynamical friction, etc.).

We can learn a lot about the universe by studying dark matter-only simulations of the universe (as presented in this chapter). Since dark matter outweighs baryonic matter by $5\times$, the gravity of the dark matter is the dominant component, and is responsible for determining when and where halos arise, how massive they are, when and how they merge together, etc. Studying such simulations therefore provides useful clues toward understanding the broader goal of the formation and evolution of galaxies. But to make further progress on that front we require an understanding of many baryonic processes. That will be the subject of Part V.

Chapter 13

Dark Matter

In this chapter we will give an overview of the observational evidence in support of dark matter, what dark matter might be, and an alternative explanation for the observational evidence.

There are many observations which demonstrate that the visible matter (baryons) and Newtonian gravity (or GR) are *inconsistent*. If one takes Newtonian gravity as correct, then the implication is that there is more matter than can be explained by the visible baryons. The other possibility – that our theory of gravity is wrong – will be discussed at the end of this chapter. See Bertone & Hooper (2018) for an excellent review of the historical developments in this area.

We have already encountered the first evidence for “missing mass”, from Zwicky (1933), who used the virial theorem to deduce that there was $> 100\times$ more mass in the Coma cluster than expected. Other early evidence of tension between the luminous matter and gravity includes Babcock (1939), who measured the rotation curve of the Andromeda galaxy (M31) and found much more mass in the outer regions than expected from the luminous matter, and Kahn & Woltjer (1959) who used the relative velocities of the Galaxy and M31 to deduce that they both contained much more mass than expected. In the 1970s, several groups, including Morton Roberts, Vera Rubin and Kent Ford, published rotation curves for dozens of galaxies and demonstrated the ubiquity of the “missing mass” in the outskirts of galaxies (see Fig. 13.1).

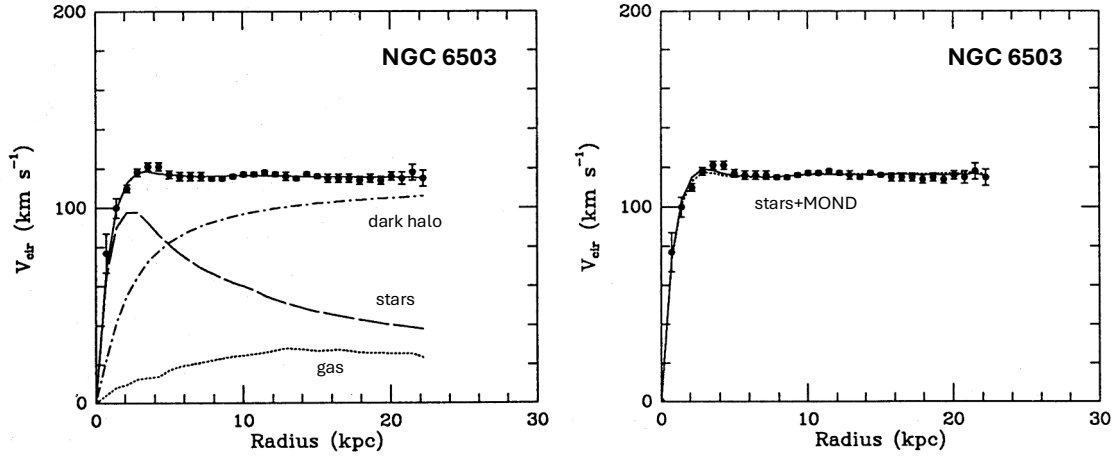


Figure 13.1: Rotation curve of the disk galaxy NGC 6503 (points). In the left panel, the data are modeled assuming a three-component model: gas, stars, and a dark halo. The resulting best-fit is the solid line. In the right panel, the data are modeled assuming only stars and modified Newtonian dynamics (MOND). Adapted from Begeman et al. (1991).

Rotation Curves measure the rotational velocity of a spiral galaxy as a function of radius. Measurement of the velocity of the gas is a particularly clean probe of the enclosed mass because gas must generally be on circular orbits. With this assumption, the circular velocity and mass are related via: $V_{\text{circ}} = \sqrt{GM/r}$. Rotation curves are therefore more robust mass estimators than the velocity dispersion profiles of elliptical galaxies because the latter requires knowledge of the velocity anisotropy to convert to a mass profile. Analysis of rotation curves does require knowledge of the inclination of the disk with respect to the observer, and a separate accounting of the gas and stellar mass contributions.

Ostriker, Peebles, & Yahil (1974) synthesized the existing observational constraints to show that $\Omega_m \approx 0.2$, i.e., that matter (baryonic and dark) contributed a substantial fraction of the critical density. It is noteworthy that the phrase “dark matter” does not appear in this work, nor in most early work on the topic. In another influential paper, Ostriker & Peebles (1974) argued that a massive halo was required in order to stabilize galactic disks. An excellent review of the state of affairs at the time is provided by Faber & Gallagher (1979).

By the 1980s it was generally agreed that the universe was full of “missing mass” that was much more extended (had a shallower density profile) than the luminous matter. The discrepancy was so large that various systematic issues (e.g., the unknown velocity anisotropy) were not seen as viable options to reconcile the theory and the data.

In the 1980s computers became powerful enough to perform N -body simulations of collisionless particles to study the implications of a universe dominated by dark matter. An influential paper by Davis, Efstathiou, Frenk, & White (1985) demonstrated that a universe dominated by cold dark matter (CDM) was able to satisfactorily reproduce many of the existing measurements of the large scale structure of the local universe (e.g., as measured by the CfA redshift survey).

As we have seen in previous chapters, the CMB offers a very powerful probe of the universe, including its overall geometry and its composition. The CMB now firmly shows that the universe is flat $\Omega = 1$ and contains a component of cold dark matter at the level of $\Omega_{\text{CDM}} = 0.25$ (i.e., a pressure-less matter component that, unlike the baryons, does not interact with photons).

One of the most dramatic observations in support of dark matter is the Bullet Cluster, which is a system of two colliding clusters of galaxies (see Fig. 13.2). The two clusters have already passed through each other once, but have not yet coalesced. The current configuration consists of two clusters of galaxies that are well-separated, while the hot gas lies in between the two clusters. The reason for this is that the gas is collisional, so as the system came together on its first approach the gas did not easily pass through. This configuration is important because the hot gas dominates the *baryonic* matter in the system. The configuration therefore offers a uniquely powerful test of the dark matter interpretation vs. an interpretation that appeals to a modification of gravity. In the former case, the “missing mass” should follow the galaxies (as both are collisionless), while in the latter case the “missing mass” will follow the gas distribution. Gravitational lensing of background galaxies provided the critical test, and it showed conclusively that the mass follows the galaxies (Clowe et al. 2006).

Any alternative explanation that does not appeal to dark matter must explain not just one or a subset of these observations, but *all* of them simultaneously. Given the very wide range of the observational constraints, that is a tall order.

The observations clearly support the notion that there is mass in addition to the baryons we can observe. But what is the missing mass?

First and foremost, the “dark” in dark matter just means that it is some non-interacting (or, more precisely, a weakly-interacting) substance/material. If it had significant interactions e.g., with radiation, or with itself, then we would expect to be able to directly observe it in some way. If it was collisional, then structure formation would proceed very differently. So

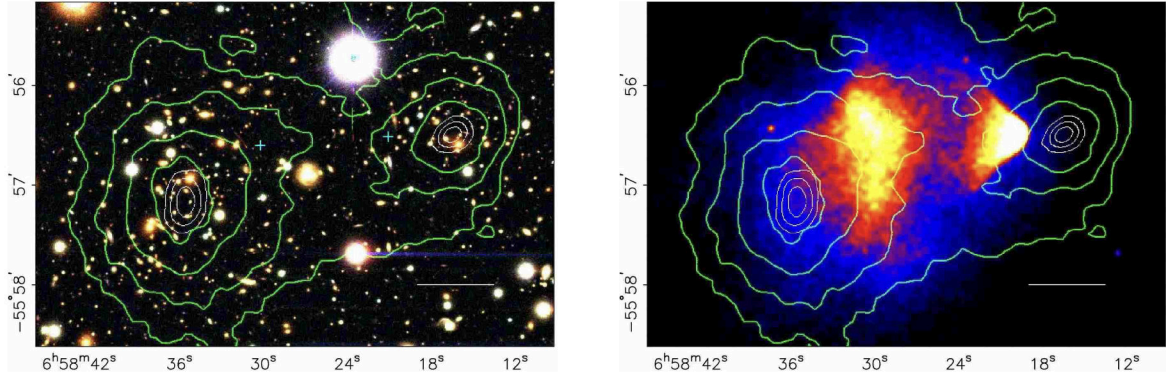


Figure 13.2: Left panel: optical image of the Bullet cluster, with weak lensing contours overlayed in green. Right panel: 500ks Chandra image, which traces the hot gas in the cluster. The hot gas outweighs the stars by $> 10\times$. The key point is that the mass contours (green) trace the collisionless stars, not the hot gas. The implication is that there must be an invisible source of collisionless matter. From Clowe et al. (2006).

our working hypothesis is that dark matter is something that interacts with gravity and the weak nuclear force.

In the early days the field was wide open as to what dark matter could be. One candidate was objects we already knew (or expected) to exist: stellar remnants (e.g., black holes, neutron stars, brown dwarfs). This class of objects is referred to as massive compact halo objects (MACHOs). Such objects would produce occasional microlensing events against a background population of stars. A major effort (the MACHO Project) was devoted to measuring this effect by using the Large Magellanic Cloud as the background set of stars. The null results from this project killed the MACHO proposal for dark matter.

The current paradigm is that dark matter is some unknown fundamental particle. There are an almost limitless number of possibilities, with particle masses ranging by some 40 orders of magnitude. The most popular candidate for the past several decades has been the WIMP, which itself is a class of models comprising a weakly-interacting massive particle. Its popularity can be traced in part to the so-called “WIMP miracle”, which we will now describe.

Most particles in the universe (dark or otherwise) were created through a process called thermal freeze out. In the case of dark matter, the basic idea is that there is some unknown particle of mass m_X . In the early universe all particles are in thermal equilibrium. As the universe expands and cools, eventually the temperature T drops below m_X (specifically $kT < m_X c^2$). At this point the dark matter particles can only be destroyed by annihilation,

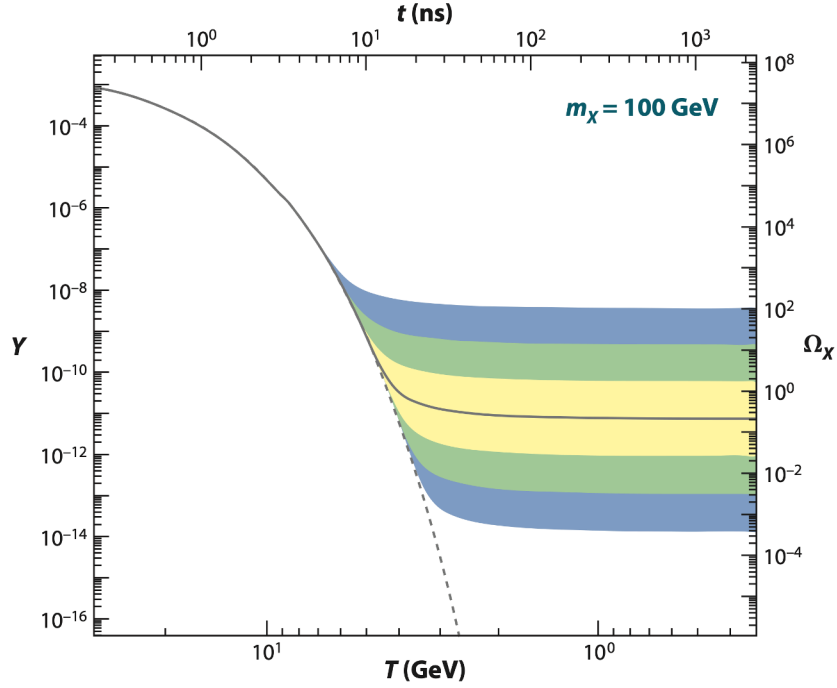


Figure 13.3: Evolution of the number density of dark matter assuming a 100 GeV particle mass. The dashed line is the result if the particle were to remain in thermal equilibrium forever (as opposed to freezing out). Shaded regions represent variation in the cross section by factors of 10, 10^2 , and 10^3 . From Feng (2010).

and their number density becomes exponentially suppressed (as $e^{-m_X c^2/kT}$). This process would drive the number density of dark matter particles to zero, except for another competing process: the continually expanding universe. Eventually the number density of particles is diluted by the expansion so that two particles cannot find each other to annihilate. This is known as “freeze out”. The particle number density approaches a constant “thermal relic density”. An example of this process is shown in Fig. 13.3.

This process is described formally by the Boltzmann equation:

$$\frac{dn}{dt} = -3Hn - \langle \sigma_X v \rangle (n^2 - n_{\text{eq}}^2) \quad (13.1)$$

where n is the number density of dark matter particle, H is the Hubble parameter, $\langle \sigma_X v \rangle$ is the thermally-averaged annihilation cross section, and n_{eq} is the dark matter number density in thermal equilibrium. On the right hand side, the first term accounts for expansion of the universe. The n^2 term arises from processes that destroy particle X (e.g., annihilation) and n_{eq}^2 arises from the reverse process that creates particles. See Feng (2010) for further discussion.

The correct solution requires numerically solving the Boltzmann equation. We can however gain insights from the following approximate solution (below we assume “natural” units in which most of the constants you’re used to are set to 1). We will define freeze out to be the epoch when $n\langle\sigma_X v\rangle = H$ (i.e., when the timescale for interaction equals the age of the universe). Recall from the Friedmann equation that a flat universe obeys:

$$H^2 = \frac{8\pi G}{3} \rho \quad (13.2)$$

and in the radiation-dominated era we have $\rho \propto T^4$. The Planck mass is $M_P = \sqrt{\hbar c/G}$, and to a reasonable degree of accuracy we can therefore write:

$$H \sim \frac{T^2}{M_P} \quad (13.3)$$

If we assume a “cold” relic, then the appropriate equilibrium number density is the non-relativistic expression:

$$n_f \sim (m_X T_f)^{3/2} e^{-m_X/T} \quad (13.4)$$

where subscripts f denote quantities at freeze out. Combining this with our freeze out condition leads to:

$$(m_X T_f)^{3/2} e^{-m_X/T} \sim \frac{T_f^2}{M_P \langle\sigma_X v\rangle} \quad (13.5)$$

Now define $x_f \equiv m_X/T_f$ and re-arrange:

$$\sqrt{x} e^{-x} \sim \frac{1}{m_X M_P \langle\sigma_X v\rangle} \quad (13.6)$$

The key insight from this expression is that values in the range $20 \lesssim x \lesssim 50$ cover a very broad range of plausible values for m_X and $\langle\sigma_X v\rangle$.

The cosmic dark matter density can be expressed as:

$$\Omega_X = \frac{m_X n_0}{\rho_c} = \frac{m_X T_0^3}{\rho_c} \frac{n_0}{T_0^3} = \frac{m_X T_0^3}{\rho_c} \frac{n_f}{T_f^3} \sim \frac{x_f T_0^3}{\rho_c M_P} \frac{1}{\langle\sigma_X v\rangle} \quad (13.7)$$

where 0 subscripts represent present-day values. We have assumed an adiabatically expanding universe, so that the entropy n/T^3 is constant.

We finally arrive at the key expression:

$$\left(\frac{\Omega_X}{0.2}\right) \sim \frac{x_f}{20} \left(\frac{3 \times 10^{-26} \text{ cm}^3 \text{ s}^{-1}}{\langle\sigma_X v\rangle}\right) \quad (13.8)$$

The “WIMP miracle” is the realization that a particle in the mass range of $\sim 100 \text{ GeV} - 1 \text{ TeV}$ will have a weak interaction cross section of $\sim 10^{-26} \text{ cm}^3 \text{ s}^{-1}$ and would therefore make

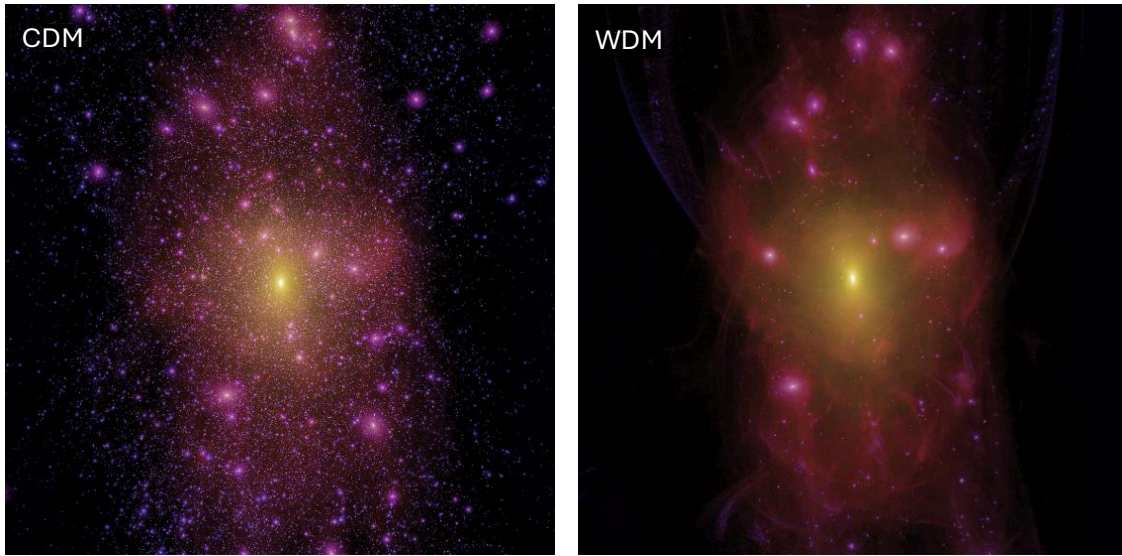


Figure 13.4: Simulation of a $10^{12} M_{\odot}$ halo in cold and warm dark matter models (left: CDM; right: WDM). Notice the large quantity of small-scale structure in CDM that is absent in WDM. This is a direct consequence of the longer free-streaming length in WDM. From Lovell et al. (2014).

an excellent dark matter candidate. Said another way, the derivation above did not make any assumptions about what type of annihilation cross section the particle must have. It just so happens that the number is at the weak scale.

This is a nice story, but as always, it is too simplistic. Direct detection experiments (e.g., from the LHC) appear to rule out particles with this cross section. Moreover, a large number of alternative models have been proposed that are just as “natural” (whatever that means) as the WIMP.

The **free-streaming length** is the distance a particle can travel before matter-radiation equality, i.e., before perturbations can grow (recall that before matter-radiation equality, the Jeans length is comparable to the horizon scale, so modes within the horizon cannot grow prior to matter-radiation equality). The comoving free-streaming length, λ_{FS} depends on the velocity of the particle prior to matter-radiation equality, and the velocity is set by the mass of the particle. Without deriving the equation, the result is (see Kolb & Turner Section 9.3.4 for details):

$$\lambda_{\text{FS}} \approx 40 \text{ Mpc} \frac{30 \text{ eV}}{m_X} \quad (13.9)$$

where m_X is the mass of the particle. All structure on scales smaller than λ_{FS} will be erased by free-streaming of the dark matter particle.

Neutrinos are very light ($\lesssim 1$ eV) particles that were highly relativistic at freeze-out. This is an example of a type of model known as hot dark matter (HDM), where the “hot” refers to the very high (relativistic) velocities of the particles at decoupling. HDM predicts very little small scale structure in the universe and has long been ruled out by observations of large-scale structure.

By contrast, cold dark matter is any particle that is non-relativistic at freeze-out. Particles with masses in the range $10 - 10^3$ GeV satisfy this (and other) constraints. CDM particles have a very small free-streaming length. Unsurprisingly, there is an intermediate case, warm dark matter (WDM), which is a class of models with particle masses in the range of ~ 1 keV. Because WDM has a lighter particle mass than CDM, it has a longer free-streaming length and hence has less small-scale structure. See Fig. 13.4 for a comparison between CDM and WDM simulations of a Milky Way-sized dark halo.

The CDM model is the currently favored model for dark matter as it not only reproduces many of the mass constraints but also broadly matches many of the large scale structure observations. However, there are many alternatives that have not been strongly ruled out, including self-interacting dark matter (SIDM), ψ DM (also known as fuzzy DM or wave DM), various flavors of WDM, etc.

Now is a good point to remind ourselves that while CDM is a very successful framework for explaining many observations, we have not yet *observed* a dark matter particle. How might we do so? The detection techniques for particle dark matter fall into three categories: 1) collider searches; 2) direct detection; 3) indirect detection.

Collider searches involve producing dark matter particles directly in a collider such as the LHC. Assuming that the particle is weakly interacting, the “detection” would be via the large amount of missing mass and momentum from the particle detector data.

Direct detection experiments rely on dark matter particles passing through detectors on Earth and physically interacting with them. As an example, a dark matter particle may interact with ordinary matter, producing a low-energy recoil in the latter. The recoil will subsequently produce scintillation radiation, which can be observed. The predicted rate depends on the unknown properties of the dark matter particle, so one is just hoping to see *something*. The challenge is that cosmic rays produce a large background, so experiments are usually performed deep underground, often in abandoned mines. Many such experiments are underway and planned. There is no convincing detection to-date.

Indirect detection of dark matter relies on the fact that particle-anti-particle annihilations will produce a range of particles, including neutrinos, γ -rays, positrons, and antiprotons. As annihilation is a two-body interaction, the rate will scale as ρ_{DM}^2 . The observed rate measured on Earth will depend not only on the DM density but also the distance from Earth and the solid angle subtended on the sky. These effects are rolled up into the “ J -factor”: $J \equiv \frac{1}{8\pi} \int dr d\Omega \rho(\mathbf{r})^2$. The Galactic Center is the region expected to have the highest annihilation signal, but the GC contains many other exotic sources (black holes, neutron stars) that can also produce high energy particles and γ -rays. For this reason the GC is a difficult location to identify dark matter annihilation. Some detections have been claimed, but they are contested. Nearby dwarf galaxies provide a cleaner potential signal, although the overall signal amplitude is lower. No detections have been claimed from dwarf galaxies.

Modified Newtonian dynamics (MOND) is a framework for understanding the discrepancy between visible matter and Newtonian gravity not as evidence for dark matter but instead as evidence of a failing of Newtonian (and GR) gravity. The key principles were laid out in a series of papers by Mordehai Milgrom in the early 1980s. The basic concept is very simple. Milgrom posited that the true gravitational acceleration, $\mathbf{a}$ is related to the Newtonian gravitational acceleration, $\mathbf{a}_{\text{N}}$ via:

$$\mathbf{a} \mu(a/a_0) = \mathbf{a}_{\text{N}} \quad (13.10)$$

where $\mu(x)$ is a function of unspecified form that has asymptotic limits of $\mu(x \ll 1) = x$ and $\mu(x \gg 1) = 1$, and a_0 is a free parameter that represents the acceleration scale at which MOND deviates from Newtonian gravity. Based on comparisons to observations, $a_0 \approx 1.2 \times 10^{-8} \text{ cm s}^{-2}$. Note that when the accelerations are large, $a/a_0 \gg 1$, MOND reduces to the usual Newtonian limit. This is by design, because Newtonian gravity is very well-tested in the limit of large accelerations (e.g., on Earth, in the Solar System). In the limit of low acceleration, the effective gravitational acceleration becomes $g = \sqrt{g_n a_0}$. If we then equate the gravitational and centripetal accelerations, $g = v^2/r$, we arrive at a very important result:

$$v^4 = GMa_0 \quad (13.11)$$

This equation has two important implications: in the limit of low accelerations, the rotation curves of galaxies will asymptote to a constant value; i.e. MOND produces flat rotation curves in the outer regions of galaxies, as observed. Second, assuming a constant mass-to-light ratio (M/L), MOND implies a relation between galaxy luminosity and asymptotic

velocity of the form $L \propto v^4$, which is the observed Tully-Fisher relation for spiral galaxies (we will discuss this relation in a later chapter). It is important to emphasize that both of these observational results were known to Milgrom, so these are not *predictions* of MOND. Nonetheless it is interesting that both observations are neatly entailed by such a simple model.

MOND has been remarkably successful at reproducing the detailed rotation curves of ~ 100 spiral galaxies, and it is this success that has been the primary driver of interest in the theory (compare the left and right panels of Fig. 13.1). For a long time MOND did not make predictions for the growth of structure, the CMB, or any relativistic phenomena (e.g., lensing). This changed in 2004, when Jacob Bekenstein proposed the first relativistic version of MOND, called TeVeS, for tensor-vector scalar gravity. Compared to GR, TeVeS contains two additional fields, three free parameters, and one free function. The implications for lensing, the CMB, etc. are still being worked out.

In spite of considerable effort, MOND has not enjoyed much success in either explaining existing phenomena (aside from rotation curves, for which it was designed) or in making new predictions. To many in the community, the Bullet Cluster result definitively ruled out MOND, although MOND-related research continues.

Part V

Galaxy Formation & Evolution

Chapter 14

Galaxies: Key Observational Facts

The first part of this course described the behavior of the homogeneous universe and the development of structure within it. We ended with the formation of gravitationally-bound, virialized dark matter objects called halos. The second part of this course will describe how the baryons within these halos cool and form stars to produce **galaxies**. This is sometimes referred to as the baryon cycle within galaxies. This field is relatively old compared to other fields within astronomy, which means that there are lots of interesting observations to understand, as well as lots of sometimes confusing jargon and terminology.

What exactly *is* a galaxy? This used to be a straightforward case of “you know it when you see it”, but with the ever-expanding observational landscape, especially in the regime of very small (low-mass) galaxies, a more precise definition is required. Willman & Strader (2012) propose the following: “A galaxy is a gravitationally bound collection of stars whose properties cannot be explained by a combination of baryons and Newton’s laws of gravity.” What this means in practice is that a galaxy is an object that resides (and presumably formed) at the center of its own dark matter halo. A challenge with this definition is that it may not apply well to galaxies that have had much of their outer dark matter halo removed by tidal stripping. Willman & Strader propose an alternative metric: the spread in metallicity of the constituent stars. It turns out that there is a clear separation between globular clusters, which show no or at most a very small spread in metallicity, and galaxies, which show much larger metallicity spreads, even at overlapping mass scales.

In this chapter we will discuss several of the most salient quantitative features of galaxies. The observational phenomenology of extragalactic astronomy is very rich; we will provide only a brief overview of the landscape.

The Milky Way (MW; also known simply as the “Galaxy”) is by far the best-studied galaxy in the universe, and in part for this reason serves as an important benchmark, both in comparison to other galaxies, and in comparison to models. The MW also has properties in common with the majority of disk galaxies in the universe. For these reasons we will discuss the properties of the MW in some detail. There are many excellent review articles focusing on the Galaxy that you may wish to consult for further details (e.g., Gilmore et al. 1989; Freeman & Bland-Hawthorn 2002; Bland-Hawthorn & Gerhard 2016).

We believe from various dynamical arguments that the MW is embedded in a dark matter halo of mass $M_h \approx 10^{12} M_\odot$ with a virial radius of ≈ 200 kpc. The total stellar mass of the MW is $M_* \approx 5 \times 10^{10} M_\odot$, which corresponds to $\approx 10^{11}$ stars.

The most prominent observational component of the MW is the disk of stars. The light profile in essentially all disk galaxies follows an exponential distribution (why this is the case is not settled). A common functional form used to fit the projected light profile, $I(R)$, is the Sersic light profile:

$$I(R) = I_e \exp \left\{ -b_n \left[\left(\frac{R}{R_e} \right)^{1/n} - 1 \right] \right\} \quad (14.1)$$

where n is the Sersic index describing the light profile shape ($n = 1$ is exponential; $n = 4$ is known as de Vaucouleurs) and R_e is the **effective radius** (also referred to generically as the scale radius) that contains half of the total light. The parameter b_n is a function that depends on n and is defined to ensure that half of the light is contained within R_e .

The scale length of the MW disk is ≈ 2.5 kpc. Note that when we refer to the “size” of a galactic component we almost always refer to the scale or effective radius. Galaxies do not have hard, well-defined “edges”.

The distribution of stars perpendicular to the disk also follows an exponential distribution. At this point in the discussion it is important to distinguish between several stellar population components within the Galactic disk, in particular, the **thin and thick disks**. The thin disk has a scale height of ≈ 300 pc, and is comprised of younger, more metal-rich and α -poor stars compared to the thick disk, which has a scale height of ≈ 900 pc. The ratio of these two components is approximately 10:1 in favor of the thin disk. The origin of this dichotomy, and whether it is really a dichotomy or more of a continuum, is the subject of much debate. Furthermore, the existence of an ancient disk of stars places strong constraints on the merger history of the MW. For example, we could not have experienced a major, disruptive merger at late times (e.g., $z < 1$) as this would have destroyed an existing old disk.

The rotation curve of the disk is flat (which was early evidence in favor of dark matter),

with a rotation velocity of $V_0 \approx 220 \text{ km s}^{-1}$ from at least $5 - 25 \text{ kpc}$. This corresponds to a period of rotation of $\approx 200 \text{ Myr}$ at the solar position (which is $R_0 = 8 \text{ kpc}$ from the Galactic center). This means that stars at 8 kpc from the center have only made 50 orbits in 10 Gyr . A flat rotation curve implies an angular velocity $\Omega \propto 1/R$, so that stars at 4 kpc from the center have made 100 rotations. This fact has important implications for our understanding of the spiral arms in our and other galaxies. If the spiral arms were fossils frozen in the disk then they would have substantially wound up within a Hubble time. Instead, we believe that spiral arms are density waves with a pattern speed much less than the disk rotation speed. But what causes the waves? Are they short or long-lived? We don't know. There is probably a fundamental difference between the origin of grand design and flocculent spirals.

The MW also contains a **bulge** in its central regions. The bulge is approximately $1/10$ of the total stellar mass, parameterized as the bulge-to-total ratio, $B/T = 0.1$. The scale-length of the bulge is $\approx 1 \text{ kpc}$. Unlike the disk, the bulge is more spherically-distributed. It now seems that the majority of the bulge is in fact contained in a long rotating bar. The origin of the bulge stars is an ongoing area of research; the favored explanation is dynamical instabilities in the disk converting disk-like orbits to more spherical orbits.

The MW also contains a stellar halo, which comprises only $\approx 1\%$ of the total stellar mass. However, due to the long relaxation and phase-mixing timescales, the stellar halo is a sensitive probe of the merger history of the MW over most of cosmic history, and for this reason is an intensively studied component of the Galaxy.

In addition to stars and dark matter, the MW also contains a large reservoir of gas. Cold gas ($T \lesssim 10^4 \text{ K}$) has settled into a rotationally-supported disk with scale height $\sim 100 \text{ pc}$. The halo contains hot gas in the range $10^5 < T < 10^6 \text{ K}$.

Essentially all disk galaxies have basic properties in common with the MW: they have a dark and a stellar halo, a disk and a bulge, etc. The differences arise in the relative mass fractions of each component, their underlying stellar populations (e.g., age and metallicity), and structural properties (e.g., scale lengths and scale heights). Indeed, much work in this field has gone into studying correlations between various properties of galaxies, as clues to their formation. Amongst disk galaxies, the most famous correlation is the **Tully-Fisher relation** (Tully & Fisher 1977), which states that the observed luminosity of a galaxy, L , is a function of the rotation velocity to the fourth power: $L \propto V_{\text{rot}}^4$.

Notation: In general discussion of star-forming and quiescent galaxies you will find the words star-forming/late-type/disk/spiral, and quiescent/early-type/elliptical used interchangeably. While galaxies generally do fall into one of these two categories, it is important to remember that morphological classification and stellar population classification are distinct, and there do exist star-forming ellipticals and quiescent disks.

We can appeal to the virial theorem to gain some insight into the origin of the TFR. Recall that $M \propto V^2 R$. If we write the mass as $M \propto \Sigma R^2$ where Σ is the surface mass density, then we have that $R \propto V^2 \Sigma^{-1}$. Now we can write the luminosity as: $L \propto M \Upsilon^{-1}$, where Υ is the dynamical mass to light ratio. Combining:

$$L \propto M \Upsilon^{-1} \propto V^2 R \Upsilon^{-1} \propto V^2 (V^2 \Sigma^{-1}) \Upsilon^{-1} \propto V^4 \Sigma^{-1} \Upsilon^{-1} \quad (14.2)$$

This basic argument recovers the TFR only if Σ and Υ are relatively constant across the galaxy population. This turns out to be a highly non-trivial result, and one that simulations struggled to reproduce for several decades.

The TFR as well as many other relations have been studied in many contexts, e.g., as functions of redshift and environment (e.g., in clusters vs. groups vs. the field). These relations serve as the fundamental points of comparison for all models of galaxy formation. We will come across more relations in later chapters.

The elliptical galaxies are distinct from the disk galaxies in many respects: the former are red, round, and are common either amongst the most massive galaxies or amongst the satellite population (galaxies in clusters are overwhelmingly dominated by ellipticals). Ellipticals also become much less common at high redshift.

There were two classic scenarios for the formation of ellipticals. In the **monolithic collapse** picture (Larson 1974) it is assumed that the stellar populations formed during the initial collapse of the galaxy ($t_{\text{SF}} < t_{\text{dyn}}$). In this scenario the ellipticals formed very early, consumed all of their gas, and evolved *passively* over most of cosmic time (passive evolution refers to evolution driven solely by the evolution of the stars within the system). In the **merger scenario** (Toomre 1977 – this is a highly recommended article to read) it is assumed that ellipticals form via the merger of two disk galaxies. This picture was informed by the spectacular images of galaxy mergers (provided for example in Arp’s 1966 Atlas of Peculiar Galaxies) as well as early N -body simulations suggesting that mergers of disks could

produce spherical systems. There is a lot of interesting history here, but the basic point is that both pictures are mostly wrong, but some elements of them survive today. What was missing in these early scenarios was any understanding of the cosmological background. The confirmation of the standard cosmological model provided a very firm backdrop on which galaxy evolution must take place. In particular, the key novelty provided by the cold dark matter cosmology is the **hierarchical** nature of the growth of structure — massive objects grow by the accumulation of a spectrum of smaller objects. You will gain experience with this hierarchical growth in HW4.

Recall from the chapters on dynamics that the distribution function, f , is a conserved quantity under the assumption of collisionless dynamics. Its observable analog, the coarse-grained DF, $\bar{f}$, can only decrease with time. We can use this fact to place constraints on the merger scenario for ellipticals (see Carlberg 1986 who first discussed this constraint). If we assume ellipticals obey an isothermal mass distribution (this turns out to be a reasonably good assumption), then the central mass density (in coordinate space) will scale as $\rho_x \propto \sigma_*^2/r_c^2$ where σ_* is the stellar velocity dispersion and r_c is a characteristic radius. The density in velocity space will scale as the width of a 3D Gaussian, i.e., $\rho_v \propto \sigma_*^{-3}$. We might therefore expect the central phase space density to scale as:

$$f_c \propto \rho_x \rho_v \sim \frac{1}{\sigma_*^3 r_c^2} \quad (14.3)$$

It turns out that ellipticals have much higher central phase space densities than spirals of the same mass. It therefore seems unlikely that one can produce most ellipticals from the mergers of present-day spirals. Historically, this was an important argument against the merger scenario. There are however two important caveats to this argument: 1) the constraint assumes collisionless dynamics, but if new stars are formed in the merger then one could increase the phase space density (though this would then predict significant quantities of young stars if the mergers were recent, which we generally do not observe); 2) the implicit assumption is that ellipticals formed from the merger of *present-day* spirals. But the mergers likely happened at higher redshifts, where the phase space density of spirals may have been higher.

As with the disk galaxies, ellipticals obey a number of scaling relations. In fact, for reasons having to do with the evolution of stellar populations (and their lack of dust), they actually obey more, and generally tighter relations than the spirals.

One of the first relations discovered was the **red sequence**, which is a tight relation in the color-magnitude diagram for quiescent galaxies. The tilt of the red sequence (in the sense that brighter galaxies are redder) was a subject of much debate 10-20 years ago, but is now

known to be mostly due to increasing metallicity in more luminous galaxies.

As we saw with the spirals, the virial theorem allows us to predict a scaling relation between the dynamical mass-to-light ratio, Υ , size, R , velocity dispersion, σ , and surface brightness, $I \equiv L/4\pi R^2$:

$$\boxed{R \propto \sigma^2 I^{-1} \Upsilon^{-1}} \quad (14.4)$$

Observationally one finds, for ellipticals:

$$\log R_e = a \log \sigma_0 + b \log I_e + c \quad (14.5)$$

with $a \approx 1.2 - 1.5$ and $b \approx -0.8$. The scatter in this relation is remarkably small, $\lesssim 0.1$ dex. The relation above is known as the **Fundamental Plane** (Djorgovski & Davis 1987). The deviation of the coefficients a and b from the predicted virial relation of 2 and -1 is known as the *tilt* of the FP. An enormous amount of work has gone into understanding the origin of the tilt, with two classes of explanations: 1) non-homology, i.e., structural variation within galaxies as a function of mass, such that there are mass-dependent coefficients relating R and I ; 2) varying Υ . Note that $\Upsilon = (M_{\text{dyn}}/M_*)(M_*/L)$, where the first term depends on the dark matter fraction and the latter depends on the stellar population. The answer turns out to be that both of these explanations 1) and 2) play some role in explaining the tilt.

The FP and its small scatter is a very strong constraint on models, in particular one must maintain the FP in the face of mergers over cosmic time.

There are two important projections of the FP:

$$L \propto \sigma^4 \quad \text{Faber – Jackson Relation} \quad (14.6)$$

$$I_e \propto R_e^{-1.3} \quad \text{Kormendy Relation} \quad (14.7)$$

The bulges of spiral galaxies are in many respects similar to elliptical galaxies in the sense that they obey the same scaling relations. This is the defining characteristic of “classical” bulges. There is another class of bulges called pseudo-bulges (of which the MW bulge is an example) that do not obey the elliptical-like relations. The distinction has to do with their presumed formation histories. The classical bulges are thought to form like the ellipticals (e.g., early star formation followed by mergers) while the pseudo-bulges are believed to form in-situ within the disk galaxy (see the review by Kormendy & Kennicutt 2004 for details).

We have mostly focused so far on the local universe ($z \sim 0$). One of the universe’s greatest gifts is the finite speed of light. This means that observing objects further away is equivalent

to viewing those objects as they were in the past. Hence, observing galaxies at higher redshifts is equivalent to observing the universe as it was in the past. Using this fact to observe the evolution of the universe (and the objects within it) is known as **lookback studies**. For example, by observing galaxies at $z \approx 0$ and $z \approx 1$ we can obtain statistical representations of the population at two epochs; these data can be used to construct or constrain models for the evolution of galaxies. This is a complementary approach to **archeological studies** in which we study the present-day properties of galaxies in detail to infer how they might have evolved. Both approaches are in widespread use, and they have complementary pros and cons (I encourage you to list a few of each).

When considering the evolution of galaxies, it can be important to distinguish between the formation epoch(s) of the stars from the assembly epoch(s) of the stars. This is especially true in a hierarchical Λ CDM universe in which there is predicted to be significant accretion and mergers over cosmic time. For example, massive elliptical galaxies generally contain very old stars (so an early formation epoch of the stars), but based on lookback studies it seems that they assembled most of their stars relatively late (e.g., $z < 1$).

Finally, one should be mindful of the many possible biases that can affect one's observations or conclusions. **Malmquist bias** (Malmquist 1922, 1925) refers to the fact that it is easier to observe intrinsically more luminous objects. Nearly every survey of stars or galaxies is flux-limited, meaning that the sample will contain an over-representation of intrinsically more luminous objects. This also means that one will observe the intrinsically more luminous objects at greater distances. If the distances are sufficiently large, then one must also contend with evolutionary corrections (objects further away will be probing a different cosmic time).

Another important effect is **progenitor bias** (van Dokkum & Franx 2001). If you are studying quiescent galaxies as a function of redshift, then you are selecting galaxies according to some criterion (e.g., by color) at each epoch. However, because galaxy color can evolve rapidly with time, at least some of the progenitors of the $z = 0$ quiescent galaxies are unlikely to be quiescent at $z = 1$. By requiring your sample to be quiescent at each epoch, you are selecting a biased subset of progenitors at each epoch. Correction for this effect can be challenging.

Chapter 15

Gas and the Road to Star Formation

Having provided a broad overview of galaxies, we are now going to spend the next several chapters diving deeper into the physics and phenomenology of various topics relevant to the formation and evolution of galaxies. In this chapter we will discuss gas in the universe. Gas (defined here as all baryons not in stars or stellar remnants) is both the fuel for the formation of stars and an important probe of the dynamics and energetics of the universe.

Understanding the distribution of gas in the universe is more complicated than the case of collisionless particles (such as dark matter or stars) because, in addition to gravity, gas can cool and heat, and is subject to collisional processes (e.g., shocks). The story in brief is this: gas is initially distributed in a manner largely tracing the overall matter. As overdensities develop and collapse to virialized objects, the gas will also collapse and virialize. One can then imagine dark matter halos containing gas that has been shock-heated to the virial temperature of the halo. This gas can then cool and form a cold interstellar medium (ISM) out of which molecular clouds and eventually stars may form. Some of this gas can be subsequently heated by star formation (or black holes) and may be ejected from the ISM or even the entire halo. In addition to virialized halo gas, there is also substantial gas in the intergalactic medium (IGM), i.e., the space in between galaxies.

A warning that this chapter contains a lot of facts on a variety of topics all under the umbrella of “gas in the universe”. It may take a few passes to absorb the information.

The universe has a present-day baryon density of $\Omega_b \approx 0.05$, i.e., 5% of the density of the universe is in baryons. We also talk about the baryon fraction as the fraction of *matter* that is baryons, i.e., $f_b \equiv \Omega_b/\Omega_m \approx 0.16$. The baryons are overwhelmingly in the form of

hydrogen and helium. We quantify the mass fraction of each component of the baryons as X for hydrogen, Y for helium, and Z for the remainder of the elements, known in astronomy as “metals”.

The metallicity plays a critical role in the thermodynamics and evolution of gas and stars even though it is always a small fraction by mass, because metals are an important coolant and source of opacity. For reference, the Sun has $Z_{\odot} = 0.015$, i.e., 1.5% of the Sun by mass is in elements heavier than He (in fact, the majority of the mass in metals is in oxygen). This value happens to be close to the peak of the metallicity distribution function (MDF) of stars in the solar neighborhood. However, most components of the baryonic universe have metallicities $\ll Z_{\odot}$. For example, the metallicity of the IGM is in the range $10^{-3} \lesssim Z_{\text{IGM}} \lesssim 10^{-2}$.

Of the baryons, we can estimate the fraction of mass in stars (including stellar remnants). It turns out this fraction is quite small: $\Omega_* \approx 0.003$; i.e., less than 10% of the baryons are in stars (Fukugita, Hogan, & Peebles 1998 provide an excellent overview of the cosmic baryon budget). We will discuss this further in the next chapter.

It is important to recognize that the gas in the universe can be found in several distinct phases. The easiest to observe and hence best-studied are the cold gas within galaxies at a temperature of $\lesssim 10^4$ K and the hot gas in groups and clusters of galaxies at temperatures of $\gtrsim 10^6$ K. A careful accounting of the mass contained in these two phases reveals that half of the baryons are “missing”. Simulations have long predicted the existence of a warm-hot intergalactic medium (WHIM) of gas at $10^5 - 10^6$ K. This phase is very difficult to observe owing to the relatively low predicted densities of the gas. This gas is predicted to trace out the large-scale sheets and filaments of the large scale structure in the universe.

Gas associated with collapsing halos is hot. How hot? If we let $kT = \mu m_p \sigma^2$, where σ is the 1D velocity dispersion and μ is the mean molecular weight (see the supplementary material at the end of this chapter; $\mu = 1$ for pure H and $\mu = 0.5$ for ionized H), then a Milky-Way halo with $\sigma = 150 \text{ km s}^{-1}$ has $T \approx 10^6$ K while a cluster halo with $\sigma = 1000 \text{ km s}^{-1}$ has $T \approx 10^8$ K.

A blackbody with temperature of 12,000 K has a peak of emission at 1 eV (this is a good set of numbers to memorize). So a MW-like halo will have a peak emission at ~ 100 eV while gas in a galaxy cluster will have a peak emission at 1 – 10 keV. The *Chandra* X-ray telescope has an imaging sensitivity in the range of $\sim 1 - 10$ keV and so is very sensitive to the hot gas in groups and clusters of galaxies.

Stars form from cold ($\lesssim 100$ K) gas. Gas therefore needs to cool from $10^6 - 10^8$ K to < 100 K in order for stars to form. How does it do this? The key concept here is the cooling function,

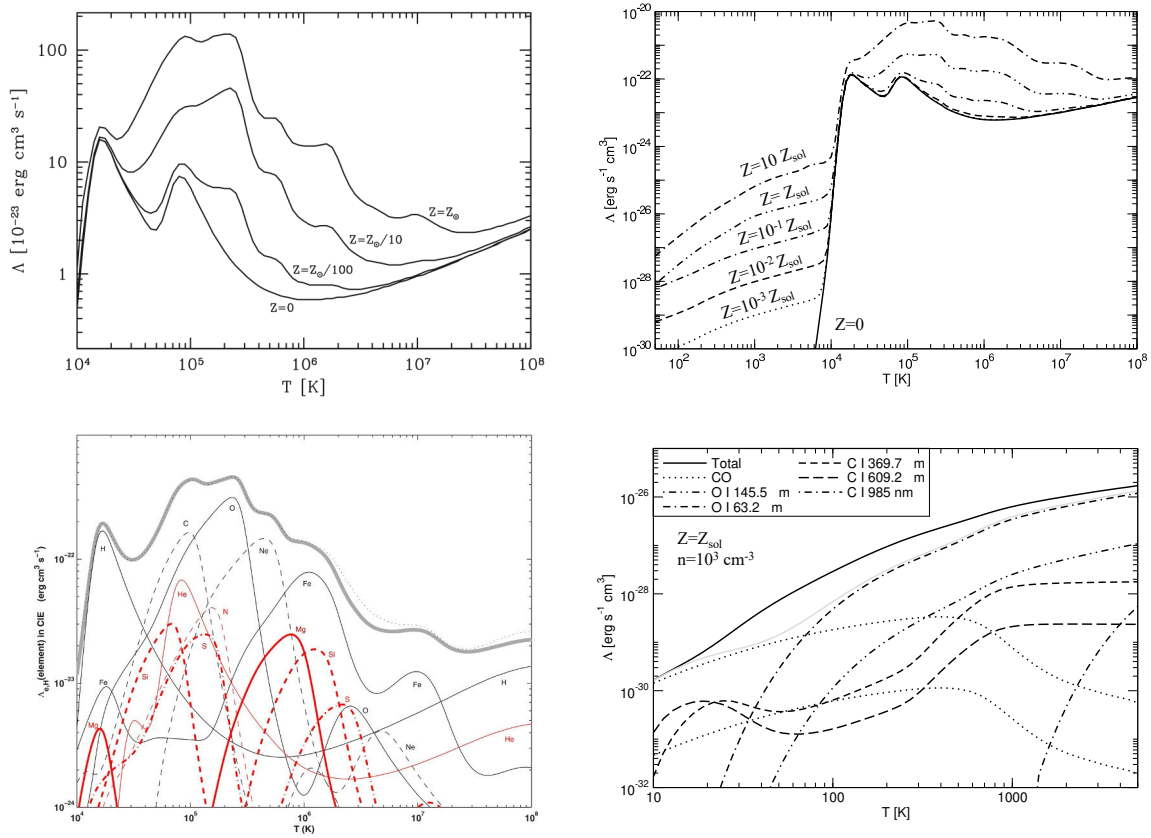


Figure 15.1: Cooling curve as a function of temperature and metallicity (top panels) and separated by component (bottom panels). Plots are from Mo, van den Bosch & White, Gnat & Ferland (2012), and Smith, Sigurdsson, and Abel (2008).

$\Lambda(T)$, which is a function that combines all of the relevant cooling mechanisms for a gas; see Fig. 15.1 for examples. In general the cooling of a gas is a two-body process (de-excitation can also cool, but collisions are required to get the excitation in the first place). The overall cooling rates are therefore frequently expressed as: $n^2\Lambda(T)$ where n is the particle number density within the gas.

The cooling time is defined as the time required to radiate away the thermal energy of the gas, i.e.,:

$$t_{\text{cool}} \propto \frac{nkT}{n^2\Lambda(T)} \quad (15.1)$$

What processes dominate the cooling function? It depends on the temperature:

- At high temperature ($\gtrsim 10^7$ K) the primary coolant is Bremsstrahlung, also known as

free-free emission. The cooling function in this regime scales as $\Lambda(T) \propto T^{1/2}$. In this regime the cooling time scales as $t_{\text{cool}} \propto T^{1/2}n^{-1}$.

- At intermediate temperatures ($10^4 - 10^7$ K) collisional excitation and ionization cooling dominate. The basic process is simple: a collision excites an atom, typically resulting in radiative de-excitation (photon emission). If the photon is able to escape the system (i.e., if the system is optically thin) then energy escapes the system, and cooling occurs. Ionization-recombination cooling can also occur, but is generally less important because the recombination rates are lower than de-excitation rates.
- At lower temperatures ($< 10^4$ K) atomic cooling is not important because collisions cannot excite atoms to the first excited state. Cooling is therefore slow below 10^4 K. In this regime molecules and dust play a dominant role; at $\lesssim 10^3$ K the cooling is due to numerous ro-vibrational energy levels available in molecules.

Notice in Fig. 15.1 that the cooling function depends strongly on metallicity at both $< 10^4$ K and in the $10^5 - 10^6$ K regime.

A key threshold is met when $t_{\text{cool}} \leq t_{\text{dyn}}$. In this regime the gas can cool efficiently and collapse. As it collapses it falls deeper into the potential well, becoming denser and cooling faster. Eventually the collapse is halted by angular momentum and the gas settles into a rotationally-supported disk with $T \sim 10^4$ K. In the opposite regime, where $t_{\text{cool}} \geq t_{\text{dyn}}$, the gas will remain in a pressure-supported hot spherical component (the hot halo). This comparison of timescales has played an important (probably too important) role in our understanding of the formation and growth of galaxies; see e.g., Rees & Ostriker (1977), White & Rees (1978). These simple arguments predicted an upper limit to the masses of galaxies of $\sim 10^{12} M_{\odot}$, which is comparable to the most massive galaxies we observe today. This agreement may be a coincidence, as lots of mass is built up by mergers, and the early picture did not include dark matter halos. Further, the time that matters is probably not t_{dyn} but the time since the last major merger, since mergers can reset the cooling time “clock” (during a major merger gas is shock-heated back to the virial temperature). See Fig. 15.2.

An optically-thin gas will tend to cool. If there is some external source of photons, e.g., from star formation, then this can heat the gas, mostly via photoionization heating. Other sources of heat are available, e.g., shocks, cosmic rays.

If heating balances cooling then we might expect an equilibrium configuration (i.e., a steady state). However, we still need to check the thermal stability of the equilibrium configuration.

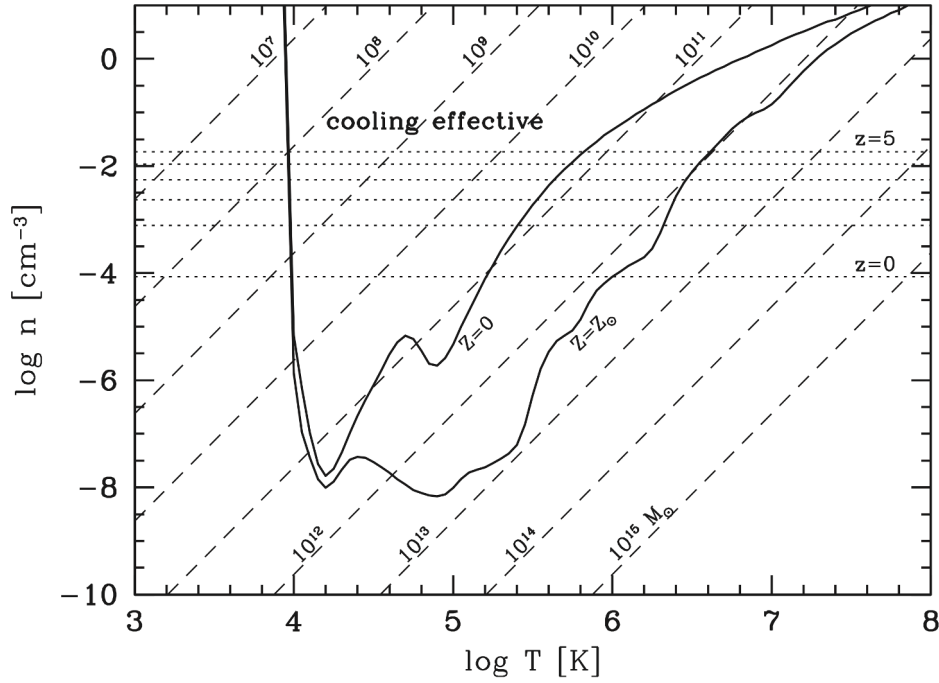


Figure 15.2: Cooling diagram in the $n - T$ plane. Solid lines show the locus of $t_{\text{cool}} = t_{\text{dyn}}$ for metal-free and solar metallicity gas. Dashed lines trace constant gas mass, while dotted lines show the gas densities for virialized halos ($\delta = 200$) at different redshifts. The calculations assume a uniform gas cloud in virial equilibrium with $f_{\text{gas}} = 0.15$. Cooling is effective ($t_{\text{cool}} < t_{\text{dyn}}$) in the region above the solid lines. The key insight from this figure is that clouds with gas masses in excess of $10^{11} - 10^{12} M_{\odot}$ (depending on metallicity) cannot cool effectively. For this reason we would expect (and observe) large reservoirs of hot gas in groups and clusters. From Mo, van den Bosch & White.

Basically, if the cooling-heating balance is perturbed, will the gas return to the equilibrium configuration or not?

We can define the loss function, $\mathcal{L} \equiv (C - H)/\rho$, where C and H are the total cooling and heating rates. $\mathcal{L}$ will be a complex function of density, temperature, metallicity, the radiation field, etc. $\mathcal{L} = 0$ corresponds to the equilibrium condition, while $\mathcal{L} > 0$ corresponds to net cooling and $\mathcal{L} < 0$ corresponds to net heating. Fig. 15.3 shows a schematic loss function as a function of temperature and density for a typical heating source. The diagonal dashed line is a line of constant pressure.

Imagine perturbations at constant pressure (isobaric perturbations). At the location P_2 , an increase in T at constant P results in net cooling, i.e., the temperature will decrease. This point is stable. At the location P_1 an increase in T will result in net heating and so T will

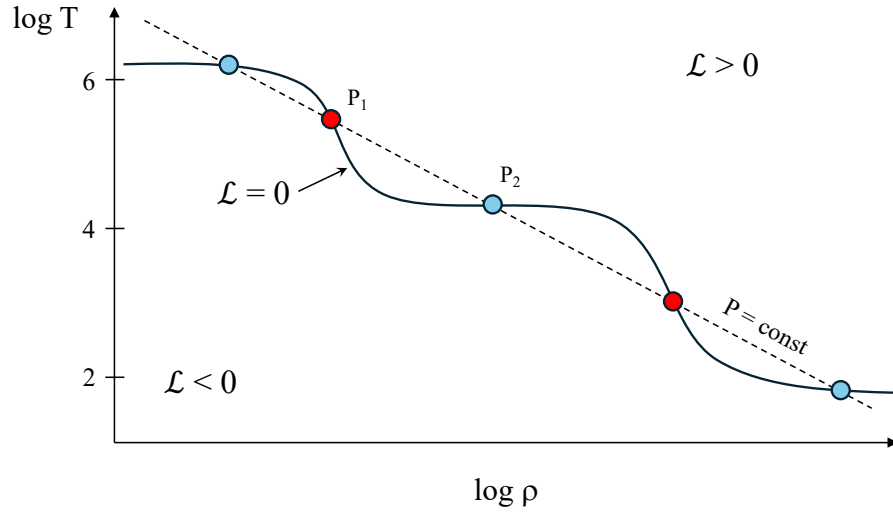


Figure 15.3: Schematic loss function as a function of temperature and density. The dashed line is a line of constant pressure. Blue circles mark points of thermal stability; red circles mark points of thermal instability.

continue to increase. This point is thermally unstable. Formally, we can see that stability requires $\frac{\partial \mathcal{L}}{\partial T}|_P > 0$. This was first worked out by George Field in 1965 (Field was also the first director of the CfA).

There are three regions of thermal stability corresponding to ≈ 100 K, 10^4 K, and 10^6 K. This is sometimes referred to as a multi-phase medium; i.e., gas in multiple phases all in pressure equilibrium. This turns out to be a good model for the gas in the Milky Way (see e.g., McKee & Ostriker 1977).

How does gas in the disk collapse to form molecular clouds? Recall our discussion of Jeans stability. One way to understand Jeans collapse is to compare the sound crossing time, $t_{\text{sound}} = R/c_s$, to the free-fall (i.e., dynamical) time, $t_{\text{ff}} = 1/\sqrt{G\rho}$ where R and ρ are the system size and density, and c_s is the speed of sound within the system. The Jeans length can be derived by setting $t_{\text{sound}} = t_{\text{ff}}$, which yields (ignoring constants of order unity):

$$\lambda_J = \left(\frac{kT}{\mu m_H G \rho} \right)^{1/2} \quad (15.2)$$

Scales smaller than the Jeans length are stabilized by pressure (sound waves), while scales larger than the Jeans length are subject to collapse (gravity wins). When a cloud collapses,

it will do so isothermally, at least initially (via cooling), so the temperature remains constant, but the density is increasing. This means that the Jeans length *decreases*, leading to fragmentation. Eventually the densities become so high that the cloud becomes optically thick and the temperature increases, ending the fragmentation process.

We are now going to expand our discussion of collapse by considering the additional stabilizing effect of the Coriolis force in a rotating disk. The full analysis is provided in MvdBW Ch 11.5.2 (see also Toomre 1964); here we provide a back-of-the-envelope derivation.

Let's consider a uniform, thin, gaseous disk rotating with angular velocity Ω , with mass surface density Σ , and with an internal pressure characterized by a sound speed c_s . Imagine a patch of gas with radius h that undergoes a small compression $h \rightarrow (1 - \alpha)h$ with $\alpha \ll 1$. The change in the gravitational force per unit mass as a result of this compression is:

$$|\Delta F_g| = \frac{G\Sigma h^2}{h^2} - \frac{G\Sigma h^2}{(1 - \alpha)^2 h^2} \approx \alpha G\Sigma \quad (15.3)$$

The pressure within the disk is approximately $p \sim \Sigma c_s^2$. Assuming that c_s^2 is constant during the compression, the resulting increase in Σ implies an increase in the pressure of order $\Delta p \sim \alpha \Sigma c_s^2$. The associated increase in pressure force per unit mass is:

$$|\Delta F_p| = \frac{\nabla p}{\Sigma} = \frac{\Delta P}{h\Sigma} \sim \frac{\alpha c_s^2}{h} \quad (15.4)$$

The increased pressure force will be sufficient to stabilize the increased gravitational force if $|\Delta F_p| > |\Delta F_g|$, which implies that the scale of the compression must satisfy:

$$h < \frac{c_s^2}{G\Sigma} \quad (15.5)$$

The centrifugal force associated with rotation is a second force capable of counterbalancing gravity. We will assume that the patch conserves angular momentum about its own center during the compression, and this angular momentum per unit mass is $j = \Omega h^2$. The centrifugal force per unit mass is $F_c \sim \Omega^2 h \sim j^2/h^3$. If we assume that j is conserved, then F_c will increase during the compression:

$$|\Delta F_c| \sim \frac{\alpha j^2}{h^3} \sim \alpha \Omega^2 h \quad (15.6)$$

For this force to stabilize the compression we require $|\Delta F_c| > |\Delta F_g|$ which implies that the scale of the compression must satisfy:

$$h > \frac{G\Sigma}{\Omega^2} \quad (15.7)$$

For all scales to be stable to compression requires:

$$\frac{c_s^2}{G\Sigma} > \frac{G\Sigma}{\Omega^2} \quad (15.8)$$

or:

$$\frac{c_s\Omega}{G\Sigma} \equiv Q > 1 \quad (15.9)$$

where Q is known as the Toomre parameter (Toomre 1964). Notice that Q is a property of the disk; if $Q > 1$ then the disk is stable to fragmentation at all scales; if $Q < 1$ then the disk is unstable, and the least stable mode (the critical mode) is $\lambda_{\text{crit}} = G\Sigma/\Omega^2$.

As always, this simple story is not the full story. It is likely that thermal instabilities trigger gravitational instabilities (e.g., a drop in temperature results in a drop in c_s and hence Q). While it is possible that this picture of fragmentation is responsible for the giant molecular clouds (GMCs), it is not even clear if GMCs are gravitationally bound objects or are transient objects forming at converging points in a larger scale turbulent flow. It is also not clear to what extent we can think of the cold ISM as a series of discrete fragmented clouds versus a hierarchical distribution of densities sculpted by turbulence and gravity. The ISM is complicated. Fig. 15.4 shows how the gas in a modern cosmological simulation distributes itself in the temperature-density plane.

Gas is everywhere, and in almost every conceivable state, from 10 K to $> 10^8$ K, and densities from 10^{-6} cm^{-3} to 10^6 cm^{-3} . How do we observe the gas? There are two basic approaches: via absorption and via emission. The former requires a brighter background source (e.g., a quasar or bright galaxy), which is an important limitation. A key advantage of absorption-based methods is that the absorption signal scales as the first power of density ($\propto n$), which means that absorption is generally sensitive to lower densities of gas than emission, as the latter generally scales as density squared (for reasons discussed earlier in this chapter). For emitting sources, one can have continuous emission due to blackbody radiation, non-thermal emission mechanisms such as free-free and synchrotron emission, as well as line emission (e.g., recombination lines). Much more details can be found in the Radiative Astrophysics and ISM courses.

With sufficient data, we can in principle fully determine the metallicity (and even abundance mixture), and pressure, density, and temperature of the gas. However, in practice the amount of data required is prohibitive, and so one often relies on a combination of measurements and physical models of the emitting (or absorbing) region to infer the physical state of the gas.

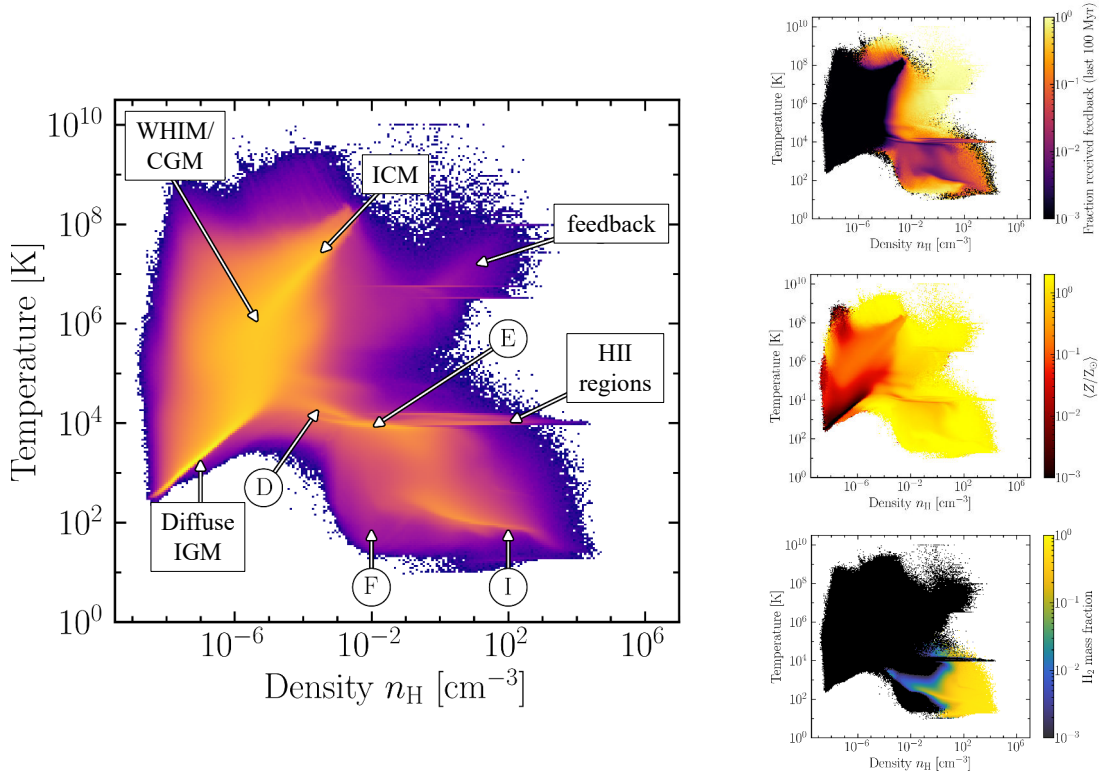


Figure 15.4: Temperature-density diagram from a cosmological simulation. Various key phases are highlighted. The left panel is weighted by the number of particles in the simulation. The right panels are weighted by the fraction of particles that experienced feedback (top), the metallicity (middle) and the H_2 mass fraction (bottom). Adapted from Schaye et al. (2025).

Owing to the very wide range in temperatures as well as both thermal and non-thermal emission mechanisms, gas can be observed at all frequencies, from the radio through X-rays.

Hot gas emission: Gas in the temperature range of $\gtrsim 10^7$ K is observable in X-rays via free-free emission. The energy-resolved spectrum allows measurement of the temperature as well as the composition of the gas. The luminosity in X-rays provides a measurement of the gas density. Hot gas in groups and clusters is routinely studied in this way. The signal is much weaker for lower-mass halos, because both the temperature and densities are lower.

Warm gas emission: Gas at temperatures of $\sim 10^4$ K can be either in the form of neutral hydrogen or ionized hydrogen, depending on the density and any ionizing sources. The neutral gas can be observed via the 21 cm H I line (caused by a spin-flip transition). Ionized hydrogen (H II) can be detected via recombination lines, including the very strong hydrogen Balmer lines in the optical wavelength range.

Cold gas emission: Gas at temperatures of $\lesssim 100$ K will mostly be in molecular form;

the most abundant molecule will be H_2 . However, H_2 has no permanent dipole, and so lacks the strong ro-vibrational transitions that make other molecules easy to detect. One must therefore rely on tracer populations, of which CO is the most widely employed because CO is relatively abundant and contains numerous transitions in the millimeter range. For example, the $J = 1 - 0$ transition is at $2.6 \text{ mm} = 115 \text{ GHz}$. Once one has determined the CO abundance (which requires also knowing the temperature of the cloud) then one adopts a CO- H_2 conversion factor, the so-called X_{CO} factor. There is an entire subfield devoted to measuring X_{CO} in different environments. Estimates of the H_2 mass come from situations in which the total cloud mass can be estimated (e.g., by assuming hydrostatic equilibrium) and then assuming that the overwhelming majority of the mass is in H_2 . Basic questions remain, such as whether or not X_{CO} varies with cloud properties. This is all very important because beyond a few kpc all we can hope to measure is the CO luminosity, and yet we want to know the total molecular gas mass.

Dust emission: In the cold ISM, where $\lesssim 100 \text{ K}$, dust is likely to form and survive. Dust is broadly defined as molecules larger than a few atoms. The thermal emission by dust peaks in the infrared (IR), and is often the dominant contribution to the IR for star-forming galaxies. With various assumptions of the dust-to-gas ratio, the dust luminosity alone can be used to estimate the cold gas mass.

The above are all emission-based probes of the gas. Absorption-based methods most commonly employ a very bright, ideally featureless source, such as a quasar (QSO), as a backlight to observe the gas along that line of sight. The most famous examples of this are the $\text{Ly}\alpha$ forest, which is a very large number (a forest) of absorption lines due to the $\text{Ly}\alpha$ hydrogen transition. The large number at different wavelengths arises from the different redshifts of the absorbers. This technique is widely used to study the IGM. More recently this technique has been used to probe the circumgalactic medium (CGM), i.e., the halo gas surrounding galaxies. The problem with measuring the absorption of a single transition is that it only measures the column density of that transition. Assumptions regarding the temperature and density are often required to derive total masses, metallicities, etc.

We can use the Milky Way as a reference point to put some concrete numbers to what we've been discussing. The mass within 15 kpc of the Galactic center is $10^{11} M_\odot$, of which half is in stars and half in dark matter, and less than 10% is in gas. Of the gas, approximately 60% is atomic, 20% is molecular, and 20% is ionized. The scale-height of the gas disk is 100 pc while the scale length is $\approx 2 \text{ kpc}$. The disk is thin. The velocity dispersion of the gaseous disk is $\sigma_g \approx 10 \text{ km s}^{-1}$, which corresponds to $\approx 10^4 \text{ K}$. The ISM in the MW is multi-phase, with all phases in approximate pressure-equilibrium. Phases include the ionized gas, warm

H I, cool H I, diffuse H₂, and dense H₂. The dust-to-gas ratio in the MW is 0.5%. This ratio is observed to correlate linearly with metallicity in other galaxies, suggesting that the dust-to-metals ratio is constant. This is not surprising, as the dust is primarily composed of metals.

Supplementary Material: Mean Molecular Weight

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 – more recent studies advocate $Z_{\odot} = 0.015$). 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 $m_{\text{H}} = 1.67 \times 10^{-24}$ g. The mass of a helium atom is (not quite exactly but close to) $m_{\text{He}} = 4 m_{\text{H}}$, and the mass of element Z_p is approximately $2Z_p m_{\text{H}}$. The average proton number of heavy elements is $Z_p = 8$, because it so happens that oxygen is the most abundant heavy element. So the mass of heavy elements can be reasonably approximated by $16 m_{\text{H}}$.

We can then estimate the number density of ions as

$$n_i = \frac{\rho}{m_{\text{H}}} \left(X + \frac{Y}{4} + \frac{Z}{16} \right) \quad (15.10)$$

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:

$$n_e = \frac{\rho}{m_{\text{H}}} \left(X + \frac{Y}{2} + \frac{Z}{2} \right) = \frac{\rho}{m_{\text{H}}} \frac{1 + X}{2} \quad (15.11)$$

(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 $n = n_i + n_e = \frac{\rho}{m_{\text{H}}} (2X + 0.75Y + 0.5Z)$.

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

$$\mu \equiv \frac{\rho}{n m_{\text{H}}} \approx \frac{1}{2X + 0.75Y + 0.5Z} \quad (15.12)$$

$$\mu_i \equiv \frac{\rho}{n_i m_{\text{H}}} = \frac{1}{X + Y/4 + Z/16} \quad (15.13)$$

$$\mu_e \equiv \frac{\rho}{n_e m_H} = \frac{2}{1 + X} \quad (15.14)$$

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).

These definitions allow us to write the gas pressure as:

$$P_i = \frac{\rho k T}{\mu_i m_H} \quad , \quad P_e = \frac{\rho k T}{\mu_e m_H} \quad , \quad P = P_e + P_i = \frac{\rho k T}{m_H} \left(\frac{1}{\mu_e} + \frac{1}{\mu_i} \right) \quad (15.15)$$

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.

Chapter 16

Star Formation

In this chapter we will discuss the observational aspects and basic theory of star formation from an extragalactic perspective. The process by which individual stars form from dense cold gas is an active area of research, both observationally and theoretically; we will gloss over many of the interesting details here in order to focus on the “global” picture of the process of star-formation on the scale of entire galaxies.

Picking up the story from the last chapter, the basic idea is that gas can become unstable (usually via cooling) and will collapse. At the smallest scales the collapse will lead to stars. When stars form they produce energy, via radiation, stellar winds, and supernovae, which provides heat to the surrounding gas. This is called feedback. The formation of stars is therefore expected to be a **self-regulating** process. As we will see below, one of the most interesting puzzles in this field is how a process that is fundamentally driven by small (e.g., AU) scales conspires to produce a remarkable degree of regularity in the galaxy population on large (e.g., kpc) scales.

As a point of reference, the Milky Way is currently forming stars at a star formation rate (SFR) of $\approx 2M_{\odot} \text{ yr}^{-1}$. The most extreme starburst galaxies have SFRs of $\sim 10^3 M_{\odot} \text{ yr}^{-1}$.

Let’s begin with the observations. How do we measure the SFR of a galaxy, or of some region within a galaxy? This turns out to be very challenging, so there are many techniques, each with pros and cons. The challenge is driven in part by stellar physics – massive stars are far more luminous than lower-mass stars (on the main sequence the mass–luminosity relation scales as $L \propto M^{\alpha}$ with $\alpha \approx 3-5$ depending on mass, and evolved stars are far more luminous than main sequence stars, further exacerbating the problem) and yet are far less common

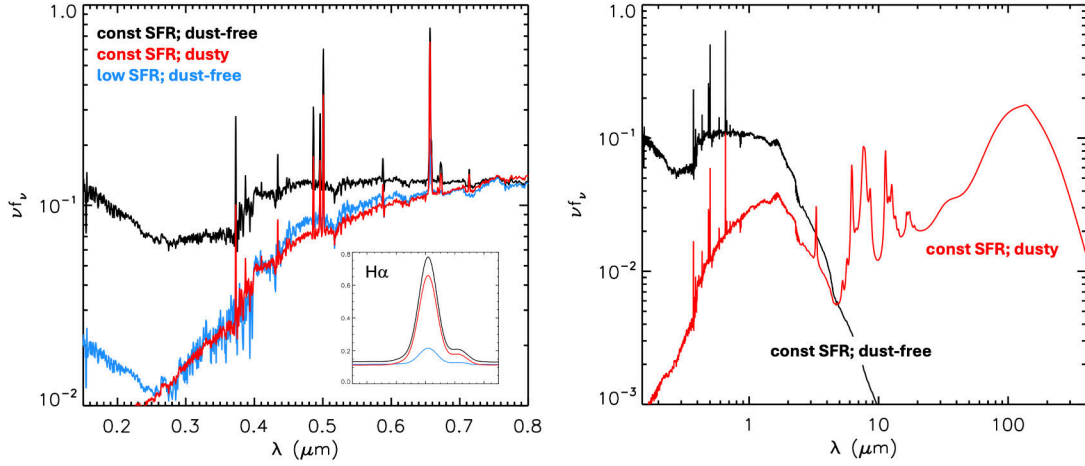


Figure 16.1: Model spectral energy distributions of galaxies. The left panel shows galaxies with a constant (high) SFR (red and black) compared to a model with a low SFR (blue). Notice the much weaker emission lines and lower UV flux in the latter. One of the constant SFR models is dust-free (black) while the other is dusty (red). Notice that the red and blue lines have similar shapes at $0.3 - 0.8 \mu\text{m}$. All models in the left panel have been normalized to the same flux at $0.75 \mu\text{m}$ in order to focus on their shapes. The right panel shows the dust-free and dusty models from the far-UV through far-IR. Notice that the dust absorbed in the UV is reprocessed and emitted in the IR.

than lower mass stars. The approach almost always taken is to use an observational proxy to count massive stars, and then extrapolate enormously to the total population using an assumed initial mass function (IMF; we will cover the IMF in more detail in later chapters, but briefly, it specifies the number of stars as a function of mass at birth; low-mass stars are observed to be far more common than massive stars).

Perhaps the most fundamental measurement of star formation comes from directly counting young stellar objects (YSOs) that are still embedded in their natal clouds. YSO lifetimes are ≈ 2 Myr; they therefore provide the best estimate of “recent” star formation. This measurement requires infrared observations and can only be done for nearby clouds in the Milky Way.

Massive stars are characterized by having high surface temperatures ($T_{\text{eff}} \gtrsim 20,000$ K) and being short-lived (lifetimes of $\lesssim 10$ Myr). Therefore, if we observe radiation indicative of a population of hot stars, we can infer that they are young and hence can estimate the SFR averaged over the past ~ 10 Myr. This approach is the most common, and most robust way of measuring SFRs in extragalactic systems.

An example of this approach is to consider hydrogen recombination lines, e.g., $H\alpha$ at $6563\text{\AA}$ and $H\beta$ at $4861\text{\AA}$. Massive stars emit lots of ionizing radiation at $> 13.6\text{ eV}$, ionizing the hydrogen surrounding them, creating H II regions. The ionized H recombines, producing a cascade of recombination photons. The whole process is fairly complicated but well-understood (public codes such as `cloudy` exist, which make realistic predictions of the resulting emission spectrum). Recombination directly to the ground state is difficult to observe because those photons quickly encounter another neutral H atom and are absorbed. Recombinations to excited states produce photons that can escape the cloud. With the aid of models, we can then directly translate the line flux in $H\alpha$ (for example) to a total ionizing flux, and from there to a total number of massive stars, and with an IMF, to a total SFR. Because the hot O stars capable of producing ionizing radiation have lifetimes of $\lesssim 10\text{ Myr}$ we often talk about $H\alpha$ as measuring the SFR on 10 Myr timescales.

Dust is always a complication. In general the effect of dust absorption increases toward shorter wavelengths. So some $H\alpha$ will be absorbed by dust, and more $H\beta$, etc. Luckily for us, quantum mechanics sets the *intrinsic* ratio of $H\alpha/H\beta$ (known as the Balmer decrement); the observed ratio therefore provides an estimate of the effect of dust on these lines. Typical corrections are ~ 1 mag of extinction – that’s a significant correction. See Fig. 16.1.

One solution to the dust problem is to consider infrared radiation, where dust absorption is less important. For example, there is a hydrogen recombination line $\text{Pa}\alpha$ at $1.8\mu\text{m}$ that is often considered a “gold standard” in SFR measurements, but it is hard to observe from the ground (JWST is starting to enable breakthroughs here).

One challenge with measuring recombination lines is that a spectrum is required, and those are observationally expensive. An alternative option is to consider ultraviolet (UV) emission at $\approx 1500 - 3000\text{\AA}$ (which only requires imaging data, and so can be easier to obtain). There are several challenges with the UV: it is even more sensitive to dust absorption, and the typical star emitting in the UV is less massive than those providing the ionizing spectrum, implying that the UV is sensitive to SFR on longer ($\sim 100\text{ Myr}$) timescales. This can be a feature if one combines UV and recombination line SFRs – one can imagine measuring the *history* of star formation by combining indicators that are sensitive to different timescales. Again, see Fig. 16.1.

Yet another option is to measure the luminosity in the mid-IR, which is sensitive to warm ($\sim 100\text{ K}$) dust. This dust is believed to be heated mostly by young stars of the same type that are dominating the UV.

Better yet is to estimate the SFR from the combined UV+IR emission. Whatever flux is

absorbed in the UV is re-radiated in the mid-IR, so by combining the two one should be less sensitive to dust absorption effects. There is a large literature on SFR indicators. There is a trade-off between more accurate indicators and more data required. In practice one uses a handful of very well-studied local galaxies to identify simple indicators that are reasonably accurate.

Another potential complication is deeply-embedded star formation, e.g., stars forming in regions with many magnitudes of dust extinction. In such cases essentially *all* of the UV through near-IR light is absorbed and re-radiated in the IR. It is very difficult to measure the SFR in such situations.

Another option is to use radio synchrotron emission. The basic physics here is that cosmic rays (CRs) are accelerated in supernovae shocks, and the accelerated CRs produce synchrotron emission when they gyrate in a magnetic field. The emission is optically thin, which is a positive feature (don't have to worry about dust). But there is no solid theory relating the synchrotron emission to the SFR, and so one must rely on empirical relations (known as the far-IR–radio correlation).

It is also worth emphasizing that many SFR indicators are a reflection of observational limitations. One example is the use of the [O II] doublet at 3727Å. This is not a great indicator for several reasons, but it is easy to observe at $z > 0.5$.

Modern techniques rely on fitting the full observed SED to models in which the SFR (and often the entire star formation history) is a free parameter. In principle, these models include all of the sensitivities and complications noted above.

Star formation is slow and inefficient. Consider the MW, which has a molecular gas mass of $M_{\text{H}_2} \approx 10^9 M_\odot$, and a free-fall time within GMCs of $\approx 10^7$ yr. One then might naively expect a SFR of $10^9 M_\odot / 10^7 \text{ yr} = 100 M_\odot \text{ yr}^{-1}$. The observed SFR is $50 - 100\times$ smaller than this. In other words, the Galaxy is converting cold gas into stars with an efficiency of a few percent per GMC free-fall time, or $\epsilon_{\text{ff}} \approx 1 - 2\%$. See Fig. 16.2.

We can express this inefficiency another way by defining the **depletion timescale**:

$$t_{\text{dep}} \equiv \frac{M_{\text{H}_2}}{\text{SFR}} \quad (16.1)$$

this is the time required to convert all of the fuel (the molecular gas) into stars at the current SFR. This value appears to be surprisingly constant for many galaxies at $t_{\text{dep}} \approx 2 \text{ Gyr}$.

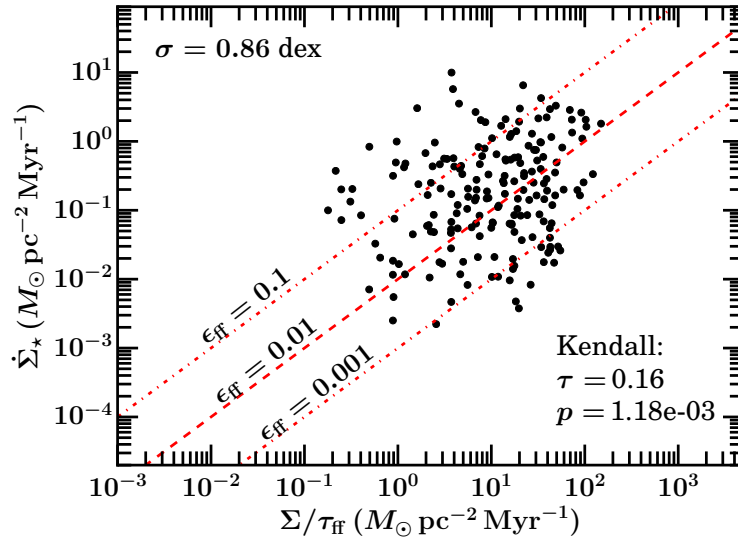


Figure 16.2: Relation between surface density of star formation and surface density of gas divided by the cloud free-fall time, for a sample of star-forming clouds in the Milky Way. Red diagonal lines indicate local star formation efficiencies of 10%, 1%, and 0.1%. From Lee et al. (2016).

There is some evidence that this depletion timescale is shorter at higher redshift ($z \gtrsim 2$) but this is controversial because one must employ the X_{CO} factor to convert observed CO emission to H_2 gas mass (see the previous chapter), and this conversion factor may not be constant.

The basic point is that our galaxy, and apparently all others, are not converting gas into stars on a free-fall time. Why? The key missing ingredient appears to be feedback, which will be the subject of a later chapter. Briefly, stellar winds, radiation pressure, and SNe are all comparable to each other and to the binding energy of the GMC. Once a few massive stars form, the energy they release will halt the collapse of the rest of the cloud. A major goal in star formation simulations is to explain this low SF efficiency.

Fig. 16.3 shows the estimated *integrated* star formation efficiency (SFE) vs. halo mass for galaxies at $z = 0$. The integrated SFE is the total stellar mass formed divided by the total available baryon mass ($M_b = f_b M_{\text{halo}}$). There are two important conclusions from this figure: 1) the peak efficiency occurs in MW-like halos (we will explore why this is the case in later chapters); 2) the peak efficiency occurs at $\approx 10\%$; i.e., star formation is not particularly efficient in an integrated sense – the universe did not figure out how to convert all or even most of the available baryons into stars. We can conclude that star formation is inefficient both locally and globally, and in both an instantaneous and time-integrated sense.

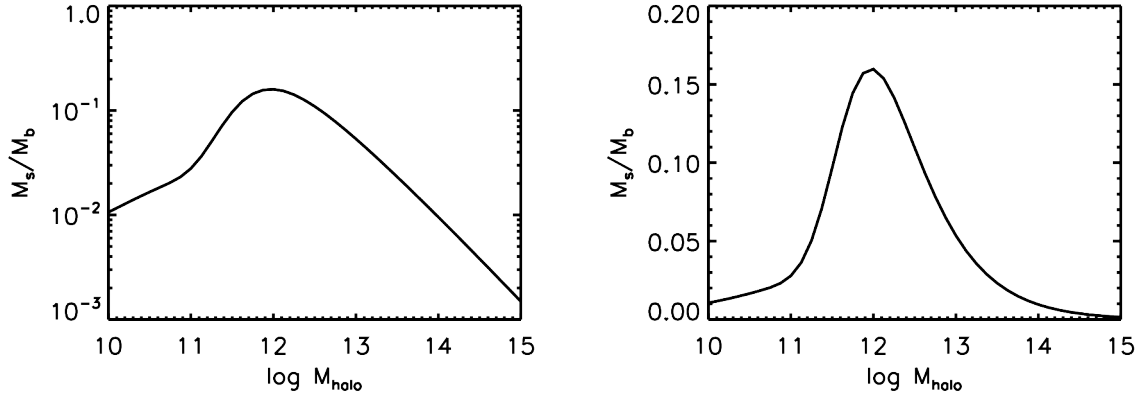


Figure 16.3: Integrated star formation efficiency — total stellar mass divided by the total available baryon mass — vs. halo mass. Left and right panels show results in log and linear axes. This relation is empirical in the sense that it was fit to a collection of observational data. It is however not *directly* measured from the data. Adapted from Behroozi et al. (2019).

We might expect a relation between the gas in a galaxy (the fuel) and its SFR. Observations, especially of galaxies outside the MW, are in projection, so the relation is often expressed as:

$$\Sigma_{\text{SF}} \propto \Sigma_{\text{gas}}^N \quad (16.2)$$

where Σ is the surface density of either SFR or gas (and is much easier to measure than the 3D density). The first attempt to constrain this relation in the MW was Schmidt (1959), and was later expanded by Kennicutt (1998), and so the relations of the form above are often referred to as Schmidt, or Kennicutt-Schmidt (KS) relations (or “laws”). Observations indicate a value for N in the range 1.4–1.5 (see Fig. 16.4). An enormous literature exists on this topic: on the observational side to measure N in different environments and on different physical scales within galaxies, and on the theoretical side to understand the physical origin of the KS relation.

If we assumed that stars are forming from gas collapsing on a free-fall time (or some fixed fraction of a free-fall time), then we might expect:

$$\rho_{\text{SFR}} \propto \frac{\rho_g}{t_{\text{ff}}} \propto \frac{\rho_g}{(G\rho_g)^{-1/2}} \propto \rho_g^{1.5} \quad (16.3)$$

where ρ_g is the gas density. It is not at all clear that stars are forming at some *fixed* fraction of a GMC free-fall time. Furthermore the above expression is in physical density, whereas

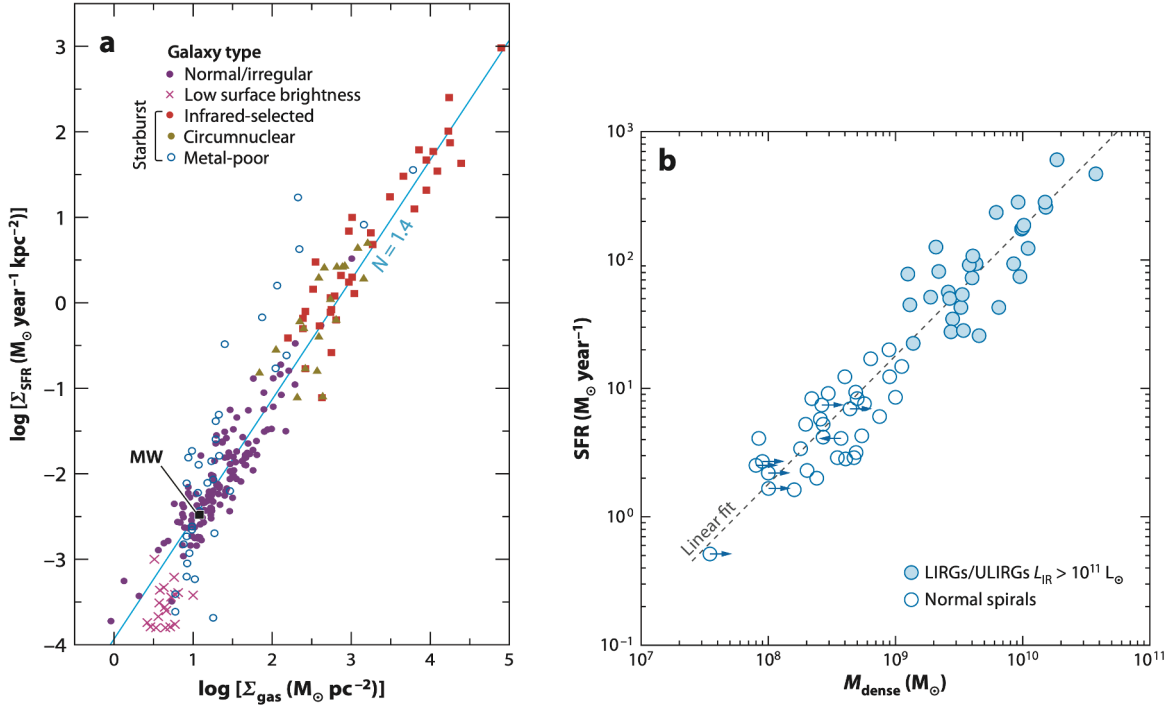


Figure 16.4: Left panel: correlation between surface density of star formation and surface density of gas (atomic and molecular). This is the Kennicutt-Schmidt relation. Right panel: Star formation rate vs. mass in cold, dense gas. This relation is linear and likely more fundamental, since stars form from dense gas. From Kennicutt & Evans (2012).

the KS relation is in projection, so one must make assumptions about the physical volume of the star-forming regions to connect the above to KS.

Until relatively recently, the KS relation made the necessary assumption that the gas density referred to the entire (cold) gas supply – both atomic and molecular. But it is the latter that is most directly connected to sites of active SF. And indeed, in more recent work that has separated the two reservoirs, one finds $\Sigma_{\text{SF}} \propto \Sigma_{\text{H}_2}$, i.e., $N = 1$ when using only the molecular gas on the right-hand side (see Fig. 16.4).

It is worth pointing out that the KS relation (whether based on total gas or molecular gas) is observed to hold over a remarkable range in SFRs, as a function of redshift, and whether one considers entire galaxies or breaks a galaxy into 100 pc regions. The relations tend to display surprisingly small levels of scatter ($\lesssim 0.3$ dex).

Two fundamental questions in this field are as follows. Why are star-forming galaxies so well-regulated? And is star formation fundamentally “top-down” (regulated by galaxy-wide properties) or “bottom-up” (regulated by GMCs or the physics of star formation)? The

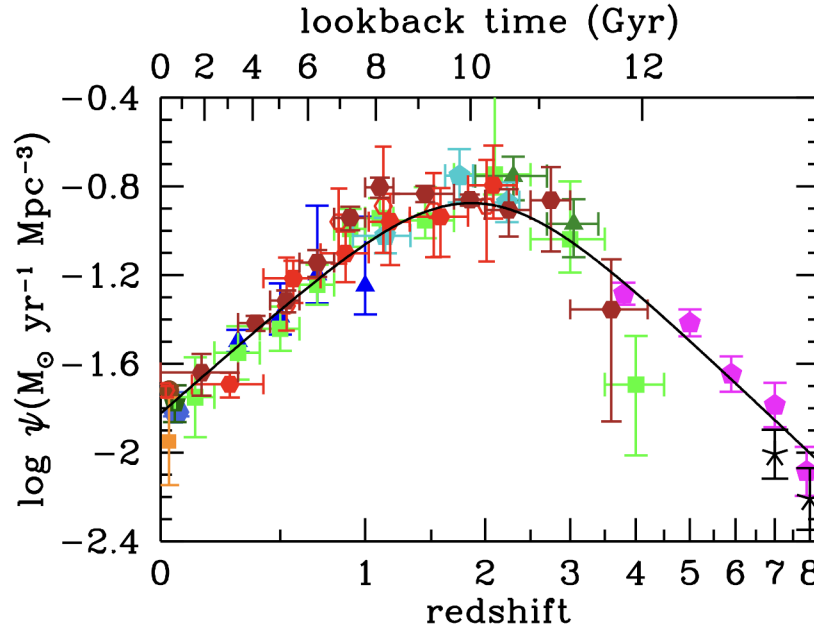


Figure 16.5: Cosmic star formation rate density over most of cosmic time. Notice that the peak of cosmic star formation occurs at $z \approx 2$. From Madau & Dickinson (2014).

latter is perhaps more intuitive but does not obviously explain the observed galaxy-scale regularity.

Fig. 16.5 shows the evolution with time of the cosmic star formation rate density, i.e., the rate of star formation in the universe per unit volume. This plot (sometimes called the Madau plot) has been widely used to test and calibrate galaxy formation models.

In this section we are going to sketch a derivation of a top-down, feedback-regulated model for star-formation. It is a useful example to see how these sorts of models are built as well as the assumptions they contain.

We start with the gas momentum equation:

$$\frac{\partial(\rho \mathbf{v})}{\partial t} = -\nabla \cdot (\rho \mathbf{v} \cdot \mathbf{v}) - \nabla P + \rho \mathbf{g} \quad (16.4)$$

where the subscript g 's referring to gas quantities have been dropped for brevity, P is the pressure, and $\mathbf{g}$ is the gravitational acceleration. If we consider a disk geometry and focus only on the z -component, then we have:

$$\frac{\partial(\rho v_z)}{\partial t} = -\nabla \cdot (\rho \mathbf{v} v_z) - \frac{dP}{dz} + \rho g_z \quad (16.5)$$

Let's now consider an average over some area A at a fixed height z :

$$\frac{\partial \langle \rho v_z \rangle}{\partial t} = -\frac{1}{A} \int_A \nabla \cdot (\rho \mathbf{v} v_z) dA - \frac{d \langle P \rangle}{dz} + \langle \rho g_z \rangle \quad (16.6)$$

where the terms in $\langle \cdot \rangle$ are averages over the area A .

The first term on the right-hand side can be separated into two components (by separating the z and xy portions of the divergence):

$$-\frac{1}{A} \int_A \nabla \cdot (\rho \mathbf{v} v_z) dA = -\frac{d \langle \rho v_z^2 \rangle}{dz} - \frac{1}{A} \int_A \nabla_{xy} \cdot (\rho \mathbf{v} v_z) dA \quad (16.7)$$

With the aid of the divergence theorem, we can re-write the second term on the right hand side as:

$$\int_A \nabla_{xy} \cdot (\rho \mathbf{v} v_z) dA = \int_{\partial A} v_z \rho \mathbf{v} \cdot \hat{\mathbf{n}} d\ell \quad (16.8)$$

this term represents the advection of momentum (ρv_z) across the edge of the area. If we assume no net flow of material within the plane of the galaxy (a reasonable assumption if we assume a galaxy in steady-state), then this term will, on average, be zero. Further, under the assumption of steady-state, the time-derivatives are zero. The gas momentum equation therefore simplifies considerably to:

$$\frac{d}{dz} \langle P + \rho v_z^2 \rangle = \langle \rho g_z \rangle \quad (16.9)$$

So far this is essentially a re-statement of hydrostatic equilibrium: pressure gradients balance gravitational forces. The key idea is now to equate the left hand side with momentum injection from stellar feedback, i.e., $\langle p/M \rangle \Sigma_{\text{SF}}$, where $\langle p/M \rangle$ is the momentum yield per unit mass of stars formed and is a property of the stellar population alone. In a plane-parallel geometry we also have $g_z = 2\pi G \Sigma_{\text{tot}}$, so we can re-cast Eq. 16.9 as:

$$\frac{d}{dz} \left(\langle p/M \rangle \Sigma_{\text{SF}} \right) = 2\pi G \rho_g \Sigma_{\text{tot}} \quad (16.10)$$

where we've added back the "g" subscripts for gas quantities. If we assume the SFR is constant across a disk scale height, h , then we can let $d/dz \sim h^{-1}$ and then we arrive at:

$$\langle p/M \rangle \Sigma_{\text{SF}} = 2\pi G \Sigma_{\text{tot}} \Sigma_g \quad (16.11)$$

where $\Sigma_g = h\rho_g$. If the gas dominates the gravity of the disk (at least at the mid-plane, where the SF is occurring), then we arrive at $\Sigma_{\text{SF}} \propto \Sigma_g^2$.

Eq. 16.11 also predicts approximately the correct normalization of the SF-gas density relation. A typical value for $\langle p/M \rangle \approx 3000 \text{ km s}^{-1}$, leading to:

$$\Sigma_{\text{SF}} = 0.1 M_{\odot} \text{ pc}^{-2} \text{ Myr}^{-1} \left(\frac{\Sigma}{100 M_{\odot} \text{ pc}^{-2}} \right)^2 \quad (16.12)$$

which is relatively close to observations, especially given the simplicity of the model.

Let's step back and consider what we've just accomplished. We have appealed to momentum injection from stars in order to provide the pressure necessary to balance the weight of the disk. This naturally leads to a relation between SF and the gas density that is a reasonably good match to the data. Remarkably, we did not appeal in any way to the physics of star formation or GMCs, or the formation of molecules, etc. Basically, this approach implies that the small-scale stuff will arrange itself in whatever way it must in order to provide the momentum from massive stars necessary to balance the weight of the disk.

The drawbacks to this approach are severalfold. First, the power-law dependence on the gas mass is incorrect (2 vs. 1.5). Second, the actual momentum coupling between stars and the gas is uncertain, and may depend on the properties of the gas. Third, there is no room for metallicity-dependence, in contrast to observations. Finally, this model explains the global inefficiency of star formation, but not the inefficiency on small (local) scales. The reader interested in understanding the bottom-up approach to star formation should consult Ch 10.2 in Mark Krumholz's excellent free textbook "Notes on Star Formation" (<https://arxiv.org/abs/1511.03457>).

The treatment of star formation in cosmological (or galaxy-wide) simulations takes several forms. The basic idea is to set a density threshold above which stars are presumed to form, often with a KS-like relation between gas density and the SFR. Some codes might also require a converging flow in the gas velocity field, and/or an extra condition on the gas temperature. Through the early 2000s, this approach was quite crude in the sense that cosmological simulations did not allow for the cooling of gas below $\sim 10^4 \text{ K}$ (owing to not tracking the formation of molecules in the simulations), and the relatively poor resolution

meant that the density threshold for SF was often very low, e.g., $n \sim 0.1 \text{ cm}^{-3}$ – note that this threshold is comparable to the *average* density of the MW ISM. Modern simulations now track the formation of molecules and gas temperatures down to $\sim 100 \text{ K}$ as well as densities $> 100 \text{ cm}^{-3}$ – much closer to the conditions relevant for GMCs.

All of these treatments of star formation are referred to as “sub-grid”, as the relevant physics for actual star formation is not captured within the simulation itself, and instead recipes are adopted to approximate the physics happening below the lowest resolution element.

Chapter 17

Stellar Populations

Galaxies are made of stars, dust, gas, and dark matter. Stars provide indispensable insight into the current state and evolutionary history of galaxies both through their kinematics and their stellar population properties. **Stellar populations** refer to the age and metallicity (or more general abundance pattern) distributions of the stars. For example, knowing the star formation history (SFH) of the stars provides unique clues to the buildup of galaxies as well as places strong constraints on models of galaxy evolution. As we will see in the next chapter, the metallicity and abundance patterns of the stars and gas provide additional clues.

The restframe UV, optical, and NIR emission of most galaxies is dominated by starlight. We can therefore hope to use UV–NIR observations of galaxies to study their stellar populations. The technique of modeling the emission of galaxies in order to infer the underlying stellar populations is known as **stellar population synthesis** (SPS). The overall structure of an SPS calculation is summarized in Fig. 17.1; the remainder of this chapter works through its ingredients in turn.

In this chapter we will be focusing on the *integrated* light from the stellar populations within a galaxy, as opposed to *resolving* the galaxy into individual stars. Obviously the latter contains far more information, but resolving individual stars is only possible for galaxies within ~ 30 Mpc (while the very faintest stars are only individually detectable within < 1 Mpc). For nearly all galaxies in the universe we observe the combined emission from millions to trillions of stars.

I encourage you to revisit the material on magnitudes, filters, and spectral energy distributions at the end of Chapter 1.

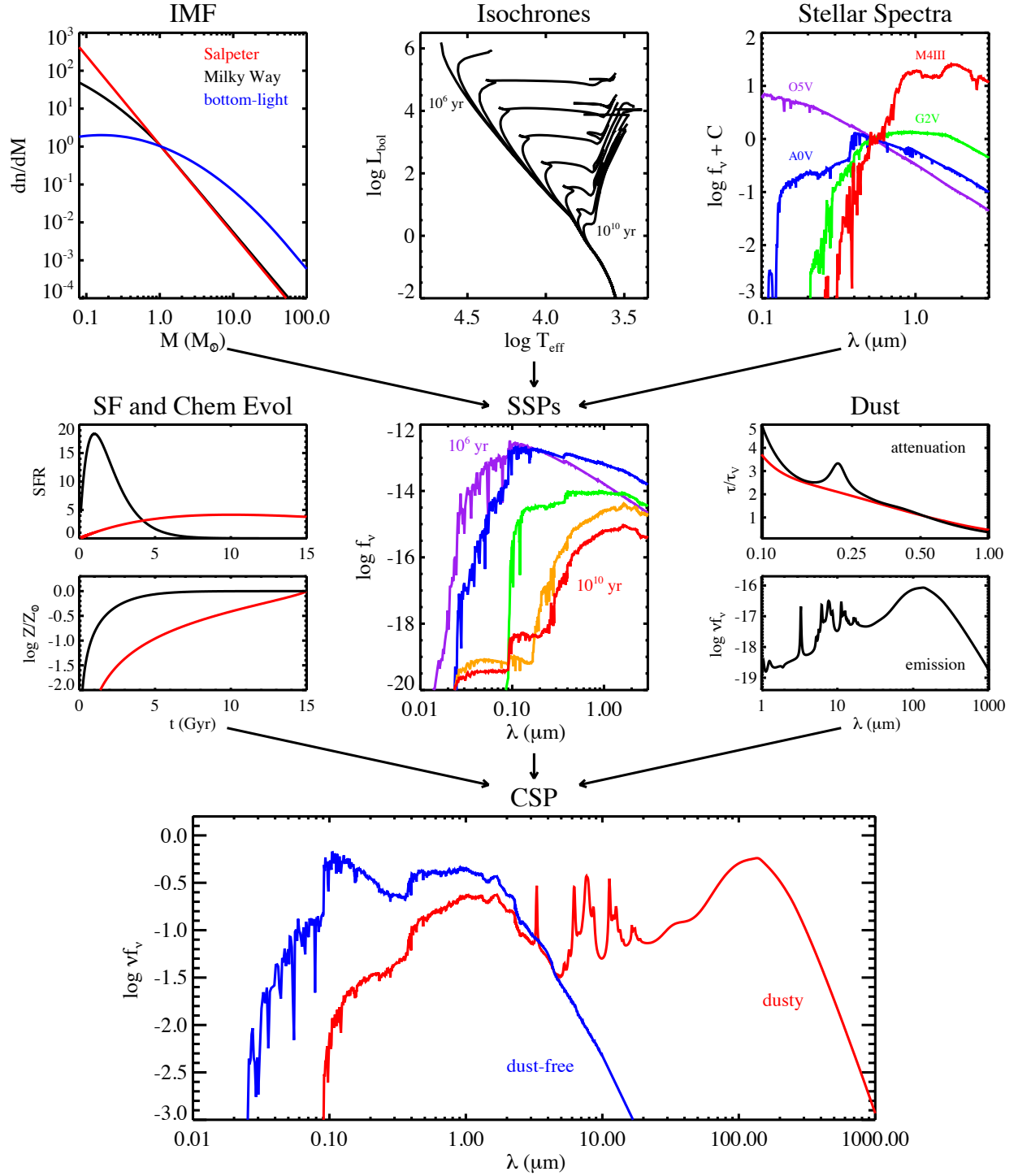


Figure 17.1: The ingredients of stellar population synthesis. The IMF, isochrones, and stellar spectra combine into SSPs, which are convolved with the SFH and MDF and attenuated by dust to yield a CSP. From Conroy (2013).

In order to build intuition for how stellar populations behave, we must begin with a brief overview of the stars themselves. The stellar initial mass function (IMF) specifies the number of stars born as a function of mass. The first empirical attempt to measure the IMF was by Salpeter (1955), who argued that the IMF was approximately a power-law: $dn/dM \propto M^{-2.3}$. Although Salpeter only measured the IMF from $\sim 0.3 - 15M_\odot$, the “Salpeter IMF” is usually taken to extend from $0.1 - 100M_\odot$. More recent determinations of the IMF in the solar neighborhood by Kroupa (2001) and Chabrier (2003) have shown that the IMF becomes shallower than the Salpeter index at lower masses. Specifically, the Kroupa IMF has $dn/dM \propto M^{-2.3}$ for $0.5 < M < 100M_\odot$ and $dn/dM \propto M^{-1.3}$ for $0.1 < M < 0.5M_\odot$. The IMF is important - it says that Sun-like stars are common while massive stars are rare.

There has been much effort to measure the IMF in other environments, but this is very challenging for several reasons. The lowest mass stars are very faint, so can only be observed at relatively close distances. Stellar models are required to convert luminosities into masses, and stars near the main sequence turnoff point require evolutionary corrections. Theory has long predicted variations in the form of the IMF with environment (e.g., metallicity and other physical conditions of the birth cloud). There is evidence based on SPS that the IMF may contain relatively more low-mass stars than the Kroupa IMF in massive elliptical galaxies, although these conclusions remain controversial.

The stellar metallicity and abundance pattern plays a key role in how stars evolve and their appearance (i.e., their spectra). We use Z to quantify the mass fraction in elements heavier than helium. For the Sun, we have $Z_\odot \approx 0.015$ although this number is perhaps surprisingly still somewhat uncertain.

There are various ways to write down the abundance of element X in a star:

$$\log \epsilon_X \equiv \log(N_X/N_H) + 12.0 \quad (17.1)$$

$$[X/H] = \log(N_X/N_H) - \log(N_X/N_H)_\odot \quad (17.2)$$

where N_X is the column density of absorbers of element X . Note that both of these definitions are number-weighted, not mass-weighted, while Z is mass-weighted.

Note that Z is a rather crude way of quantifying the metallicity of a star or population, as the *mixture* of elements contained within “ Z ” can vary from star to star. In the simplest case we adopt the assumption of *solar-scaled mixtures*, in which the relative abundance of the elements comprising Z is fixed to the solar composition. As we will see in the next

chapter, allowing the mixture to vary makes the problem both more complicated and more interesting.

Stars occupy well-defined locations in the HR diagram ($L-T_{\text{eff}}$), including the main sequence, red giant branch, red clump (horizontal branch), etc.¹ On the main sequence there is a relation between a stars luminosity and its mass, which is approximately $L \propto M^4$ – the exact power-law index varies with mass. We can estimate the main sequence lifetime via:

$$t \propto \frac{E}{L} \propto \frac{M}{M^4} \propto M^{-3} \quad (17.3)$$

scaling this to the lifetime of the Sun:

$$t_{\text{MS}} = 10 \text{ Gyr} (M/M_{\odot})^{-3} \quad (17.4)$$

Now, if we adopt $L \propto M^4$ and a Salpeter-like IMF of $\phi(m) = dn/dM \propto M^{-2}$ then we find that the integrated luminosity:

$$L_{\text{tot}} = \int_{M_{\text{lo}}}^{M_{\text{up}}} L(M) \phi(M) dM \approx M_{\text{up}}^3 \propto t^{-1} \quad (17.5)$$

where M_{up} is the upper limit of the main sequence (which is evolving toward lower masses with time due to the mass-dependent lifetime of the main sequence).

We can draw a few conclusions regarding the nature of simple (coeval; i.e., born at the same time) stellar populations under the assumption of a Salpeter IMF and emission dominated by the main sequence (i.e., Eq. 17.5). 1) At any time, the luminosity is dominated by the most massive stars still living; 2) the luminosity will decrease rapidly as time progresses, owing to the rapid evolution of the main sequence turnoff mass. In principle, a measurement of the spectrum from a galaxy tells us the properties of the most massive stars still living, which, through stellar evolution, allows us to estimate the age of the population (e.g., the presence of O stars tells us that the population is $< 10^7$ yrs old).

The complications to this simple account are severalfold: 1) most galaxies are composite populations of stars with a range of ages (and metallicities); 2) we have ignored the effect of post-main sequence evolution, which is a significant complication. Although post-MS

¹If you have never taken a course on stellar evolution and these terms are unfamiliar to you, then I encourage you to read up on the subject so that you have some basic familiarity – it will be important for this and the next chapter, and for HW5.

evolution is short in time relative to the main sequence ($\approx 10\%$ of the MS lifetime), post-MS phases can be very bright compared to the MS (the RGB is $10 - 10^3$ times more luminous than the main sequence). The proper treatment requires use of tables of stellar evolution calculations covering all relevant phases of stellar evolution for all stellar masses. This will be the subject of HW5.

The effect of dust on the observed properties of galaxies is a significant complication that must be addressed. Dust has several effects on the emergent properties of galaxies: it absorbs starlight, it scatters starlight, and it reprocesses starlight by absorption and mostly thermal remission in the infrared. An important concept is the **extinction curve**, which quantifies the extinction as a function of wavelength. The extinction curve depends on the properties and relative abundance of the dust grains (e.g., carbonaceous vs. silicate, and the grain size distribution). The fact that there are relatively more small grains than large grains implies that extinction is greater at shorter wavelengths.

One will sometimes see the words extinction and attenuation used interchangeably, but there is an important distinction between the two. Extinction measures the loss of light along a given line of sight, e.g., to a single star within the Galaxy. One measures the extinction by measuring the observed flux and dividing by a model for the intrinsic flux of the source. Extinction is the combined effect of true absorption and scattering. Attenuation refers to the effect of dust on an entire galaxy, e.g., by considering the observed emission from a galaxy compared to the intrinsic emission from the starlight alone. In the case of attenuation, there are many stars being affected by dust in various ways, as they each have a different column density of dust in between the source and the observer. Furthermore, starlight that is scattered out of a particular line of sight can be scattered back into the line of sight of the observer. To treat this problem correctly requires radiative transfer modeling.

There are various ways of quantifying the effect of dust and the extinction curve. A_V is the amount of absorption (in magnitudes) in the V -band, and is often used to normalize the extinction curve. The color excess, $E(B - V)$ is the difference between the observed and intrinsic $B - V$ color, and is defined as $E(B - V) \equiv A_B - A_V$. The ratio of the total to selective extinction provides a measure of the slope of the extinction curve, and is defined as $R_V \equiv A_V / E(B - V)$. In the Milky Way there is a relation between R and the column density of gas. The canonical value for the MW is $R_V = 3.1$. We might expect, and observations confirm, that R_V varies not only within our Galaxy but also from galaxy to galaxy.

If we define the ratio of observed to intrinsic flux as $F_{o,\lambda} / F_{i,\lambda} = e^{-\tau_\lambda}$ where τ_λ is the optical

depth, then $A_\lambda \approx 1.086\tau_\lambda$.

In stellar population synthesis one will often invoke a two-component dust model, in which stars younger than $\sim 10^7$ yr experience attenuation due to their birth cloud (the giant molecular cloud), and all stars experience dust associated with the diffuse ISM.

Having now assembled the key ingredients (the IMF, stellar evolution and dust), we can now consider stellar populations. These are the upper and middle rows of Fig. 17.1; we now work down to the bottom of that figure.

The first key ingredient is the simple stellar population (SSP), which combines an IMF, stellar evolution, and stellar spectra to predict the evolution in time of a coeval, mono-metallic population. Specifically, an SSP is a prediction of the time-evolving spectrum of a single burst of star formation, i.e., stars born all at the same time from a single homogeneous cloud of gas. The observational analog of SSPs are the open and globular clusters.

SSPs can be computed via

$$L_{\text{SSP}}(t) = \int_{M_{\text{lo}}}^{M_{\text{up}}(t)} L(M; t) \phi(M) dM \quad (17.6)$$

which is a deceptively simple expression. The luminosity-mass relation, $L(M)$, is given by a combination of stellar evolution and bolometric corrections (or spectral libraries) and in practice involves lots of lookup tables and interpolations.

Composite stellar populations (CSPs) involve the combination of many SSPs convolved by the star formation history (SFH) and metallicity distribution function (MDF) of the system, along with the effect of dust:

$$L_{\text{CSP}}(t) = \int \int \text{MDF}(Z) \text{SFH}(t - t') L_{\text{SSP}}(t', Z) e^{-\tau(t')} dt' dZ \quad (17.7)$$

where t' is the age of the population, and t is the age of the universe. Note that the above equation has made the simplifying assumption that the SFH and MDF are independent of one another, but the equation can be easily generalized.

We have a situation in which L_{CSP} is a function of the MDF, SFH, and dust. The overarching goal is to measure the latter quantities given observations of the SED and/or the spectrum of a galaxy. In practice, one makes large grids of models as a function of SFH, etc. and then

compares those grids of models to observations (e.g., of photometry or spectra). This turns out to be a hard problem, and is the subject of a large effort in the community.

Beatrice Tinsley was one of the pioneers in this field and wrote a remarkable review in 1980 that is still worth reading today (more recent reviews can be found in Walcher et al. 2011 and Conroy 2013).

The field of galaxy evolution has been deeply affected by stellar population synthesis. Prior to approximately 2004, most observational results were expressed in terms of direct observables, e.g., colors, luminosities, and distributions thereof. Following the publication of widely available, reasonably accurate SPS models (in particular Bruzual & Charlot 2003), researchers began using the models to translate observables into physical quantities, such as total stellar masses, star formation rates, etc. Nowadays it is common to hear talk of “observed” stellar mass functions, but it is important to remember that these are highly model-dependent statements.

Chapter 18

Chemical Evolution

The element abundances in both stars and gas are affected by many processes: stellar evolution and death (nucleosynthesis), the flow of matter into and out of galaxies (the baryon cycle), and the system star formation history. Measurement of the element abundances therefore offers the promise to learn about these inter-related physical processes. At the simplest level we refer to the “metallicity”, Z , as the mass fraction in elements heavier than helium. It is however more interesting and informative to consider the detailed element abundance patterns and how those elements may evolve in a particular system over time. This is the domain of a field called *chemical evolution* (the phrase derives from the fact that the elements in the periodic table are sometimes referred to as “chemical elements”, i.e., the elemental ingredients of chemistry), and is the subject of this chapter. Our goal is to highlight the primary physical processes shaping the chemical evolution of galaxies.

Chemical evolution describes how elements that are produced by stars are ejected into the ISM and how that material is then mixed into the ISM and incorporated into future generations of stars, or in some cases ejected from the galaxy into the intergalactic medium (IGM).

In general, chemical evolution is a function of several important processes and timescales:

- The stellar initial mass function (IMF), which sets the fraction of stars that are massive and will evolve and die, compared to the long-lived stars that effectively serve as “sinks” when considering the evolution of elements (i.e., elements incorporated into low-mass stars with lifetimes greater than the age of the universe will never again be available for chemical enrichment of the galaxy).

- Stellar yields, which specify the amount of mass in a given element ejected from a star upon its death. These are set by stellar evolution theory and models for supernovae explosions and can be quite uncertain depending on the element.
- The star formation history (SFH), which in combination with stellar lifetimes sets which stars are dying and returning material to the ISM as a function of time.
- Gas inflows and outflows into and out of the galaxy.
- The timescale(s) for the ejected material to mix with the ISM and become available for a subsequent generation of star formation.

Each of these ingredients represents a significant uncertainty. The first two items above are set by small-scale star formation processes and stellar evolution, while the latter three can be (somewhat) reliably predicted by hydrodynamic simulations of galaxies. Chemical evolution, probably more than any subfield in astronomy, relies strongly on many different subfields (stellar evolution, supernovae explosions, galaxy evolution, gas dynamics, etc.), making it both a very rich and complicated subject.

For the most part, the composition of the stellar atmosphere is unchanging during the lifetime of a star (we know this both from theory and observations of stars in open and globular clusters). Exceptions to this include Wolf-Rayet stars (these are very massive stars that have lost most of their outer envelopes due to winds), certain main sequence turnoff stars where diffusion is important, and stars on the RGB (in particular, for the elements Li, C, and N). Aside from these well-studied exceptions, the observed photospheric abundances can be used as a reliable proxy for the composition of the gas cloud from which the star was born.

It is important to note that gas-phase and stellar abundances offer different views of the chemical history of the galaxy. The gas-phase is an approximately instantaneous measure of the composition within the galaxy, as the gas responds quickly (on a mixing timescale) to injection of new material. The stars, being long-lived, offer a more complicated view of the chemical evolution. The element abundances of the longest-lived stars (e.g., $\lesssim 1M_{\odot}$) in particular offer an integrated view of the chemical evolution over the history of the galaxy. This is exciting, as it means we can use observations of present-day stars to learn about the history of chemical evolution in a galaxy.

We are now going to sketch out a basic **chemical evolution model**. The equations are straightforward but we have to define many variables.

Let's start with some basic governing equations:

$$M = g + s \quad (18.1)$$

where M is the total mass in the system, g and s are the mass in gas and stars (and stellar remnants). All of these variables will in general be functions of time.

The change in time of the total mass is governed by gaseous inflows, F , and gaseous outflows, E (both of which are rates and so have units of $M_\odot \text{ yr}^{-1}$):

$$\frac{dM}{dt} = F - E \quad (18.2)$$

The rate of change of the gas is given by inflows and outflows, as well as the formation of stars (SFR; ψ), and the rate of ejection of mass from stars, e :

$$\frac{dg}{dt} = F - E + e - \psi \quad (18.3)$$

It follows from these definitions of terms that the rate of change of the mass in stars is given by:

$$\frac{ds}{dt} = \psi - e \quad (18.4)$$

All we have done so far is to translate the basic physical processes into variables and simple equations.

The quantity e is computed according to:

$$e(t) = \int_{m_{\tau=t}}^{M_{\text{up}}} (m - m_{\text{rem}}) \psi(t - \tau(m)) \phi(m) dm \quad (18.5)$$

where ϕ is the IMF, $\tau(m)$ is the stellar lifetime of a star of mass m , and m_{rem} is the remnant mass left behind, which is a function of the initial mass m . The initial-final mass relation (IFMR) is the subject of much theoretical and observational work. For example, we believe that stars with $m > 40M_\odot$ leave behind black holes (the IFMR in this regime is essentially unknown); stars with $8.5 \lesssim m \lesssim 40M_\odot$ leave behind $1.4M_\odot$ neutron stars; and stars with $m \lesssim 8.5M_\odot$ leave behind white dwarfs with an empirical IFMR of $m_{\text{WD}} \approx 0.48M_\odot + 0.077m$. For a Salpeter-like IMF, these IFMRs imply that most stellar remnant mass in the universe is in the white dwarfs.

We can write down the governing equation for the evolution of metal mass in the gas in a similar fashion as above:

$$\frac{d(gZ)}{dt} = e_Z - Z\psi + Z_F F - Z_E E \quad (18.6)$$

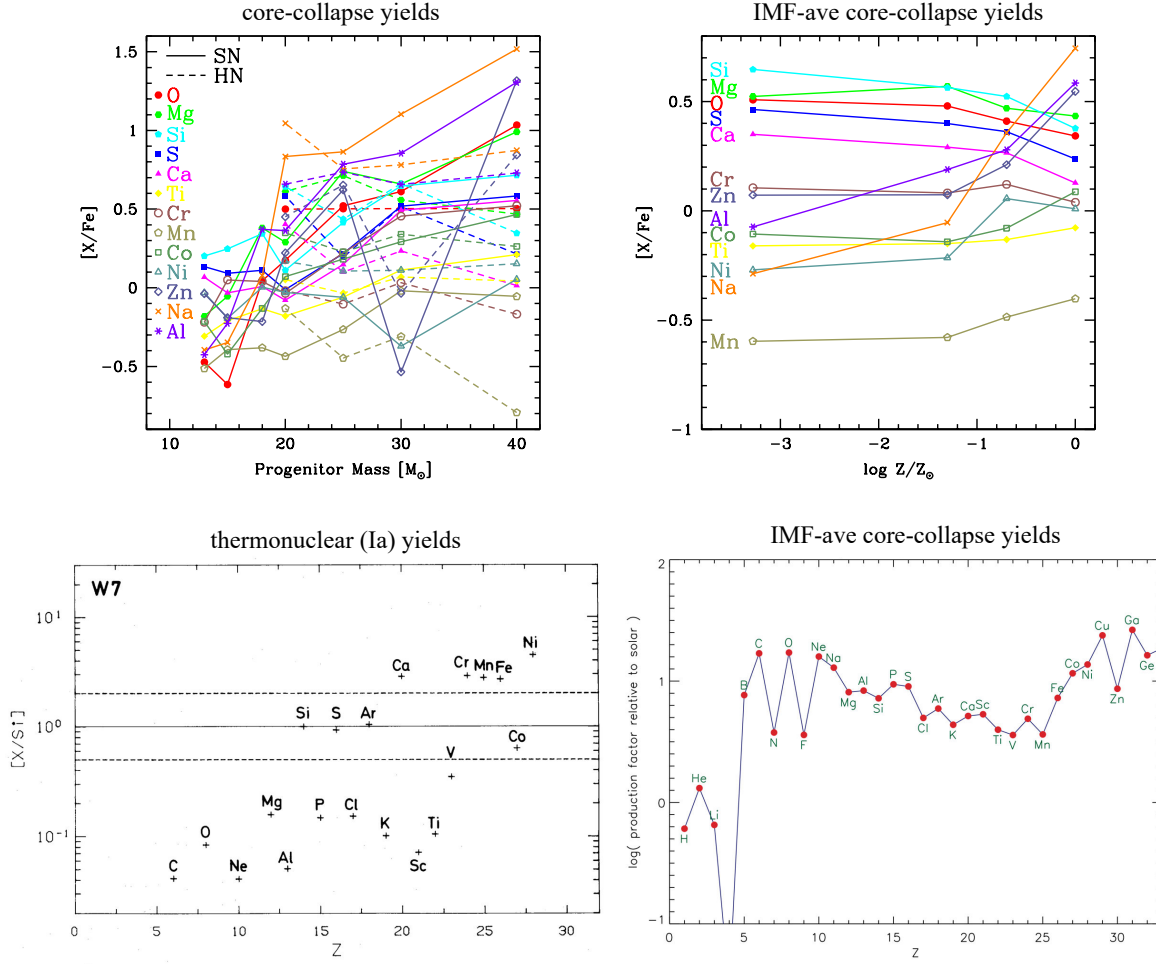


Figure 18.1: Theoretical nucleosynthetic yields. *Top left:* Core-collapse yields at solar metallicity as a function of progenitor mass. Yields are shown both for supernovae (SN) and hypernovae (HN) models. From Kobayashi et al. (2006). *Top right:* Core-collapse yields averaged over the IMF, shown as a function of metallicity. From Kobayashi et al. (2006). *Bottom left:* Yields for the thermonuclear explosion of a white dwarf. From Nomoto et al. (1984). *Bottom right:* IMF-averaged yields for core collapse supernovae at solar metallicity as a function of element charge. From Woosley et al. (2002). The top right and bottom right display similar information at solar metallicity and are the results of two different groups' efforts at computing yields.

Recall that Z is the mass fraction in metals, so that gZ is the total mass in metals in the gas phase. The first term on the RHS, e_Z , represents the total rate of mass in metals ejected from stars, the second is the metal mass loss due to star formation, the third is addition of metal mass due to inflows, and the last is loss due to outflows. Z_F and Z_E represent the metallicity of the inflowing and outflowing material, respectively, and both can in principle

be different from Z . A homogeneous wind is one in which $Z_E = Z$ and a metal-enhanced wind is one in which $Z_E > Z$. There is much discussion in the field as to whether winds (outflows) are metal-enhanced or not. If the wind is being driven by supernovae, whose ejecta will be metal-enhanced, then one can imagine the wind being metal-enhanced. In principle the metallicity of the wind is measurable, but it is not easy to do so.

The term e_Z hides a lot of detail. We can express this term as:

$$e_Z = e_Z^{\text{CC}} + e_Z^{\text{Ia}} + e_Z^{\text{AGB}} + e_Z^{\text{other}} \quad (18.7)$$

where the term has been separated into contributions from the three main nucleosynthetic sites: core-collapse supernovae, thermonuclear explosions of white dwarfs (Type Ia SNe) and AGB stars, as well as the possibility for other sources, e.g., neutron star mergers, which may be relevant for certain very heavy elements. Each of these terms has a different time dependence, which as we will see below is an important feature. Fig. 18.5 summarizes which sites are believed to dominate the production of each element in the solar system, and Fig. 18.1 shows examples of the computed yields.

The **stellar yield** is an important concept in chemical evolution. It is the mass in an element i that is *freshly produced and ejected* from a star. The stellar yield for a given element will be a function of the stellar initial mass and composition (through the effects of stellar evolution). The yield may be zero either because no new mass is synthesized (e.g., the Sun will never produce new iron, so the solar yield for iron will be zero), or because the material is synthesized but none of it is ejected into the ISM. What is returned to the ISM is what matters for chemical evolution, so if the material remains locked in a stellar remnant (as for example much of the carbon and oxygen remains locked in a white dwarf) then from the perspective of chemical evolution the material might as well never have been synthesized at all. The yield may also be negative, as is the case e.g., for Li, which is destroyed at high temperatures.

If we define $S(t)$ as all the stars ever born up to time t , i.e.,

$$S(t) = \int \psi(t') dt' \quad (18.8)$$

then we can define α as the fraction of mass remaining in long-lived stars and stellar remnants, so that $s(t) = \alpha S(t)$. We can then define the return fraction, R as $R = 1 - \alpha$. For a Kroupa (Galactic) IMF, for a population of stars with an age of 10 Gyr, we have $\alpha = 0.6$ and $R = 0.4$; i.e., 40% of the initial formed mass has been returned to the ISM in the form of stellar winds or supernovae ejecta. For a Salpeter IMF these numbers are closer to 0.75 and 0.25.

We can now define the **true yield** as the mass of an element i freshly produced and ejected by a generation of stars in units of the mass that remains locked in long-lived stars and remnants (the reason for this somewhat confusing definition will become clear in a moment):

$$p_i = \alpha^{-1} \int_{m_\tau}^{m_{\text{up}}} p'_i(m) \phi(m) dm \quad (18.9)$$

where $p'_i(m)$ is the stellar yield for element i of an individual star.

We can make considerable progress in solving these equations by invoking three simplifying assumptions:

1. The instantaneous recycling approximation assumes that stars return metals to the gas phase instantly after they are born. This is a reasonable assumption for massive stars, which have short lives, but is a bad assumption for lower mass stars and Type Ia supernovae, as we will see later.
2. The instantaneous mixing approximation assumes that fresh ejecta from stars is instantly mixed with the existing ISM and available for incorporation into new generations of stars. The validity of this assumption is questionable and depends on the detailed hydrodynamics of the system in question. It's probably not a terrible assumption in most cases. An example of where this assumption would break down is if matter ejected from stars is incorporated into the hot gaseous halo, which has a long cooling time. In such a case the metals ejected would not be available for incorporation into later generations of stars until the hot gas cools.
3. A single-zone model. The assumption here is that the system can be treated as a single physical “zone”, i.e., no spatial information. In reality we observe metallicity gradients within most large galaxies, which tells us that this assumption is dubious.

Using the above definitions and the assumption of instantaneous recycling we can re-write Eq. 18.6 as:

$$\frac{d(gZ)}{dS} = \alpha p + RZ - Z - Z_E \frac{E}{\psi} + Z_F \frac{F}{\psi} \quad (18.10)$$

where we have substituted $dS = \psi dt$ and divided both sides by ψ . In Eq. 18.6 the single parameter e_Z represented the total amount of metals ejected from stars; we have replaced this term with two terms above: αp and RZ ; the former is just the true yield renormalized to the total initial mass formed (the amount of freshly synthesized metals), and the latter is the amount of original metals in the star “recycled” back to the ISM.

We can set $RZ - Z = -\alpha Z$ and if we divide Eq. 18.10 by α we have:

$$\frac{d(gZ)}{ds} = p - Z \left(1 + \frac{E}{\alpha\psi} \right) + Z_F \frac{F}{\alpha\psi} \quad (18.11)$$

where the assumption of instantaneous mixing has allowed us to set $Z_E = Z$.

We will now spend some time describing a simple but very important and useful model for chemical evolution, the closed box model (sometimes known as the ‘simple model’). The basic assumptions of the closed box model are those we listed above (instantaneous recycling and mixing, and single-zone), as well as no inflows and no outflows, i.e., $E = F = 0$ (hence the name closed box). This model has as initial conditions $g(0) = M$, $Z(0) = 0$ and $s(0) = 0$ where M is the mass of the system (which is constant in the closed box model).

Our previous equations simplify considerably:

$$dg = -ds \quad (18.12)$$

and Eq. 18.11 becomes:

$$\frac{d(gZ)}{ds} = p - Z \quad (18.13)$$

applying the chain rule to the LHS leads to:

$$Z \frac{dg}{ds} + g \frac{dZ}{ds} = p - Z \quad (18.14)$$

The first derivative is just -1 , so that the $-Z$ terms on the LHS and RHS cancel, leaving us with:

$$g \frac{dZ}{ds} = p = -g \frac{dZ}{dg} = -\frac{dZ}{d \ln g} \quad (18.15)$$

Integrating:

$$\int_M^g -p d \ln g' = \int_0^Z dZ' \quad (18.16)$$

results in:

$$\boxed{Z = p \ln(\mu^{-1})} \quad (18.17)$$

where $\mu \equiv g/M$ is the gas mass fraction. Eq. 18.17 is the metallicity of the gas and freshly formed stars as a function of the gas fraction. Note that time has dropped out of the equations – the gas mass fraction sets the “clock” in the evolution. By now it should be clear why we defined the true yield, p , in such a convoluted way.

We can define the effective yield via:

$$y_{\text{eff}} \equiv \frac{Z}{\ln(\mu^{-1})} \quad (18.18)$$

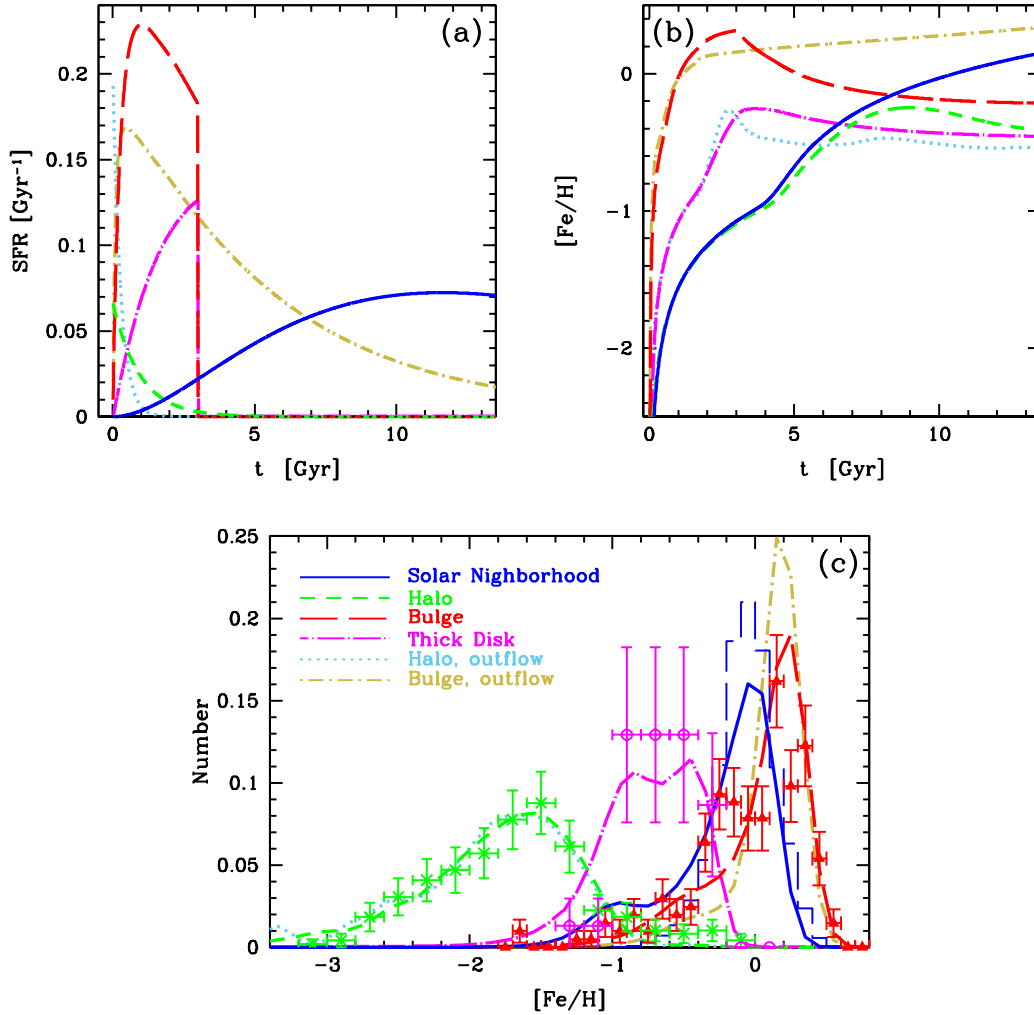


Figure 18.2: Result of a chemical evolution model tuned to reproduce the observed metallicity distribution functions (MDFs) for the bulge, thick disk, solar neighborhood, and stellar halo of the Milky Way. The top panels show the implied SFR and metallicity vs. time, while the bottom panel shows the resulting MDFs. Each component has different parameters for the star formation timescale and the degree of inflow and outflow. From Kobayashi et al. (2020).

this might seem redundant with Eq. 18.17, but note that the effective yield is defined in terms of observable quantities, the gas-phase metallicity and the gas fraction. Comparison of y_{eff} and the true yield, p , tells us to what extent a system deviates from the closed box model. Low-mass galaxies deviate substantially from the closed-box prediction (the true yield is $p \approx 0.01$). We will explore possible explanations below.

We can use the above expressions to compute the distribution of metallicities of the stars,

known as the metallicity distribution function (MDF). Starting from Eq. 18.17 we have:

$$\frac{g}{M} = 1 - \frac{s}{M} = e^{-Z/p} \quad (18.19)$$

$$\frac{s(Z)}{M} = 1 - e^{-Z/p} \quad (18.20)$$

so that:

$$\frac{ds}{dZ} \propto e^{-Z/p} \quad (18.21)$$

or:

$$\boxed{\frac{ds}{d\ln Z} \propto Z e^{-Z/p}} \quad (18.22)$$

This is the closed box MDF. On a log-log plot the MDF is linear at low metallicity with an exponential cutoff at the true yield.

There is a historical puzzle called the G-dwarf problem. Observations show many fewer low-metallicity stars than predicted by the closed box model. This was originally detected in G-dwarf stars, hence the name, but the problem persists in all stellar types. There were many proposed solutions, including problems with data, changing the IMF, or a breakdown in the assumptions of the closed box model. This last proposal is likely correct, and in particular it is now believed that significant inflow of gas during the main buildup phase of the system is the key to reproducing the observations. In our modern cosmological paradigm this arises naturally and can be easily understood, but it was something of a puzzle in the 1970s when the basic cosmological framework was still in flux. Fig. 18.2 shows the result of a modern chemical evolution model that includes inflows and outflows.

Let's return to the assumption of instantaneous recycling. As stated earlier, this is a good assumption for elements produced primarily in massive stars, but it is a bad assumption for elements produced in large quantities by AGB stars and Type Ia supernovae. Table 18.1 lists the relevant timescales and elements. Unfortunately, analytic chemical evolution models cannot be computed without the assumption of instantaneous recycling, so we are left to perform numerical integrations.

For Type Ia SNe the key quantity is the delay time distribution (DTD), which quantifies the rate of Type Ia SNe per year per solar mass of material formed, as a function of the age of the population:

$$\text{DTD} \approx \left(\frac{0.1 \text{ Gyr}}{t} \right) \times 10^{-12} \text{ SNe yr}^{-1} \text{ M}_{\odot}^{-1} \quad (18.23)$$

for $t > 10^8$ yr. The DTD has been estimated from observations but the starting time is not well-constrained.

Table 18.1. Main Sources of Element Production

type	timescale	elements produced
massive stars ($> 10M_{\odot}$)	1 – 50 Myr	α -elements (O, Mg, Si, Ti) and Fe-peak elements
AGB stars ($1 - 10M_{\odot}$)	> 50 Myr	CNO, s-process elements (e.g., Ba, Sr)
Type Ia SNe	> 100 Myr	Fe-peak elements (V, Cr, Mn, Fe, Co, Ni)

By far the most important observational diagnostic for star formation timescales inferred from element abundances is a plot of $[\alpha/\text{Fe}]$ vs. $[\text{Fe}/\text{H}]$ (referred to as the Tinsley diagram), where the latter is a ratio of generic α elements (often O or Mg are used as proxies) to Fe. The basic story is the following: at early times, e.g., before the first Ia SNe explode, the $[\alpha/\text{Fe}]$ ratio is high due to core collapse SNe. If stars are formed rapidly enough, then the iron production from CC SNe can be sufficient to drive the metallicity close to the solar value ($[\text{Fe}/\text{H}] \sim 0$). However, if star formation proceeds slowly, then the ejecta from Ia SNe will have time to be incorporated into subsequent generations of stars; this will drive the

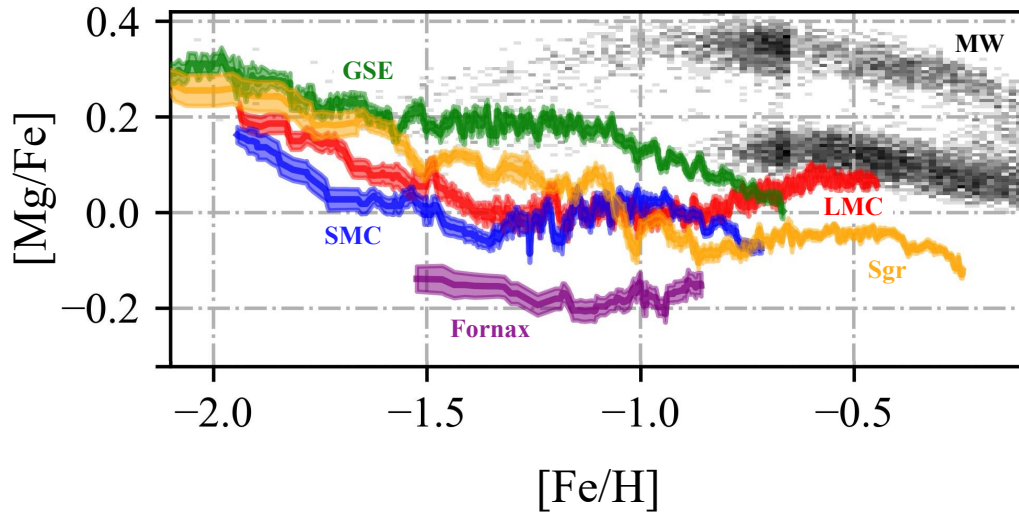


Figure 18.3: Tinsley diagram showing $[\text{Mg}/\text{Fe}]$ vs. $[\text{Fe}/\text{H}]$ for stars from the Milky Way (MW), the GSE accreted dwarf that now comprises the bulk of the stellar halo, as well as the Sagittarius (Sgr), Fornax, and LMC and SMC dwarf galaxies. The MW displays a clear bimodality, while the other systems do not and so are shown as rolling averages in $[\text{Fe}/\text{H}]$. From Hasselquist et al. (2021).

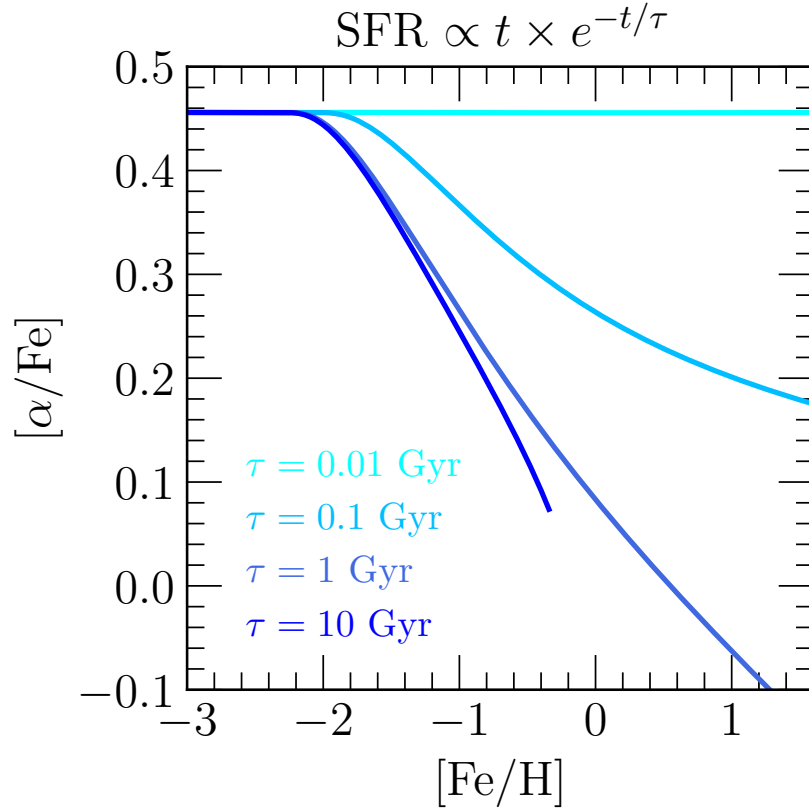


Figure 18.4: $[\alpha/\text{Fe}]$ vs. $[\text{Fe}/\text{H}]$ from a chemical evolution model for different star formation histories. Courtesy of James Johnson.

$[\alpha/\text{Fe}]$ ratio toward lower values. Exactly when Ia SNe start contributing elements relative to how far along the system is in consuming its gas supply will determine what trajectory the system will make in the $[\alpha/\text{Fe}]$ vs. $[\text{Fe}/\text{H}]$ plane. Fig. 18.4 shows these trajectories for a range of star formation histories.

One sees significant structure in this diagram when considering different stellar populations. For example, the MW disk stars have a drop in $[\alpha/\text{Fe}]$ at relatively low metallicity, while the MW bulge drops off at higher metallicity. Different galaxies also occupy different sequences in this diagram (see Fig. 18.3), indicating different star formation histories and inflow/outflow histories. A fairly universal phenomenon is the elevated $[\alpha/\text{Fe}]$ ratio for the most metal-poor stars. These stars clearly form out of pure Type II SNe ejecta, with little/no contribution from other nucleosynthetic sites. Perhaps most remarkably, the *average* star in massive elliptical galaxies is α -enhanced. One explanation for this is that the SF timescale in such systems is very short ($< 10^9$ yr). But this is an outstanding question.

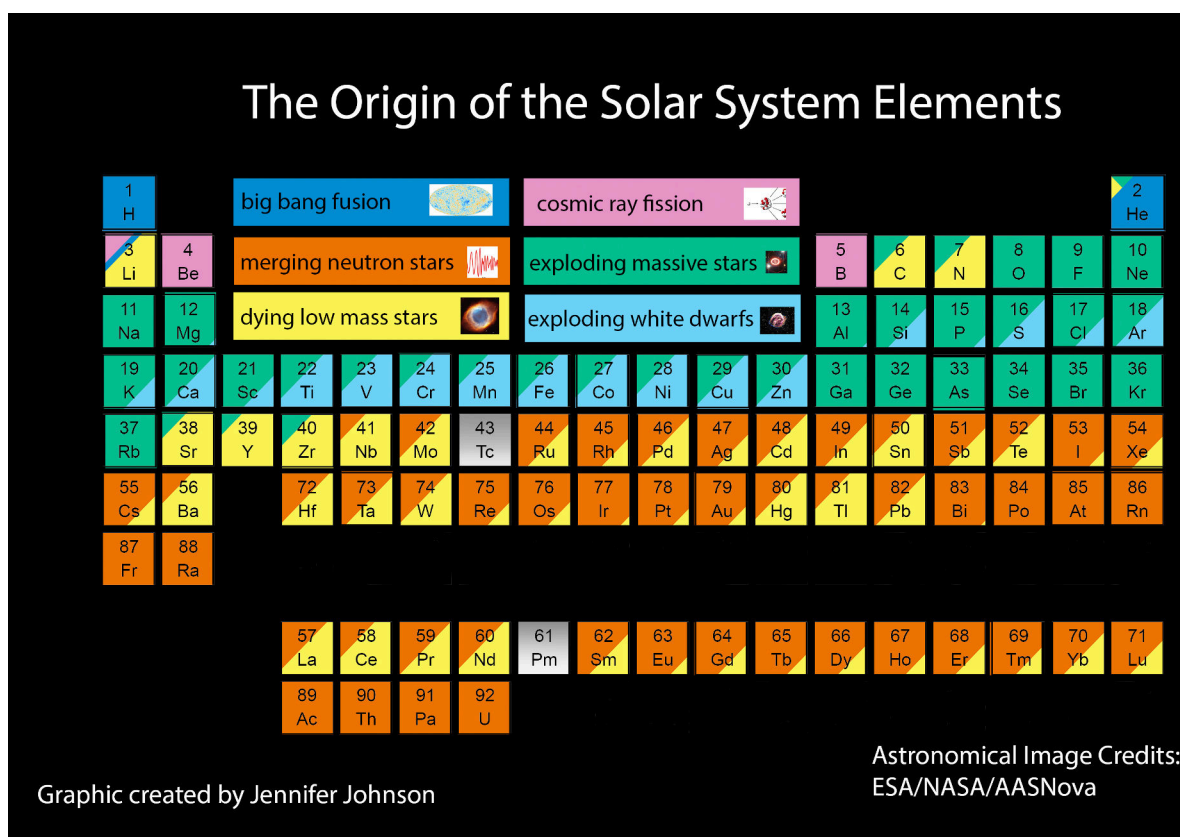


Figure 18.5: Schematic representation of the nucleosynthetic origin of the elements in the solar system. The relative importance of various channels will vary across the Galaxy and across cosmic time, and the detailed breakdown is subject to ongoing debate. Courtesy of Jennifer Johnson.

Chapter 19

Stellar & Black Hole Feedback

So far, we have covered the state of gas in the universe, how that gas can cool to form stars, and some basics of the star formation process. As we discussed previously, star formation is both locally and globally very *inefficient*, which means that only a small fraction of the cold gas is forming stars at any given time. The reason for this inefficiency is largely a consequence of feedback, which describes the basic process by which the formation and/or growth of a population (stars, black holes) produces energy and/or momentum that in turn slows or quenches the growth of that population. Feedback is believed to be the key process that gives rise to galaxy-wide scaling relations such as the Tully-Fisher and Faber-Jackson relations ($L \propto V^4$ and $L \propto \sigma^4$, respectively), and the observed scaling relations between galaxy mass and black hole mass, e.g., $m_{\text{BH}} \propto \sigma^{4-5}$ (known as the M- σ relation).

Let's begin with some basic arguments. The energy produced by a typical supernova (SN) is $\epsilon_{\text{SN}} \approx 10^{51}$ erg, and we expect 1 SN per $100M_{\odot}$ of stellar mass formed (assuming a Kroupa IMF, and ignoring Type Ia SNe, which explode on much longer timescales than is relevant for regulating star formation). A giant molecular cloud (GMC) has a typical mass of $10^5 M_{\odot}$ and typical radius of 50 pc. The binding energy of such a cloud is $E_b \sim GM^2/R \sim 10^{49}$ erg. We conclude that 1 SN can easily unbind a GMC. However, it turns out that SNe are too “late” to efficiently regulate star formation in GMCs because of a mis-match in timescales. The free-fall time of a GMC is 1 – 10 Myr, while the time it takes for a massive star to evolve and explode is 5 – 10 Myr. This is an important clue that some other source of energy must be available to provide the necessary growth-regulating feedback at the GMC scale. Fig. 19.1 shows several examples of stellar feedback in action.

We can make a similar argument on a galaxy-wide scale. The Milky Way galaxy has a halo



Figure 19.1: Examples of stellar feedback in action. *Left:* The M82 dwarf galaxy, showing large-scale outflows traced by $H\alpha$ emission (red). *Top right:* the 30 Doradus star-forming region in the LMC. A large concentration of new stars has created a cavity within the larger scale distribution of gas. *Bottom right:* Infrared image of the Carina nebula. A large population of bright O stars is located North of the image and is in the process of heating and removing the surrounding gas. Image credits: NASA, ESA, and The Hubble Heritage Team (STScI/AURA); NASA, ESA, CSA, STScI, Webb ERO Production Team; NASA/JPL-Caltech/N. Smith.

mass and size of $M_h \approx 10^{12} M_\odot$, $R_{\text{vir}} \approx 200$ kpc, and a gas mass of $M_g \approx 10^{10} M_\odot$, so the binding energy of the gas is $E_b \sim GM_h M_g / R_{\text{vir}} \sim 5 \times 10^{57}$ erg. If we imagine that only 10% of the available SNe energy efficiently couples to the surrounding gas, then we require $\sim 10^8$ SNe, or $10^{10} M_\odot$ of newly formed stars. The relevant timescale is the Galaxy’s dynamical time, $t_{\text{dyn}} \approx 10^8$ yr (and is the fastest timescale on which a system can form stars), so the implied SFR necessary to unbind the galaxy is $10^{10} M_\odot / t_{\text{dyn}} \approx 100 M_\odot \text{ yr}^{-1}$. Recall that the current SFR of the Milky Way is $1 - 3 M_\odot \text{ yr}^{-1}$.

We will consider here two classes of feedback: energy-driven and momentum-driven. If $t_{\text{cool}} < t_{\text{dyn}}$ then the thermal energy will be radiated away before it can do anything useful (it would be as if the energy was never there to begin with). However, momentum cannot be radiated away, so it is available to do “work” over a longer period of time. Therefore, in

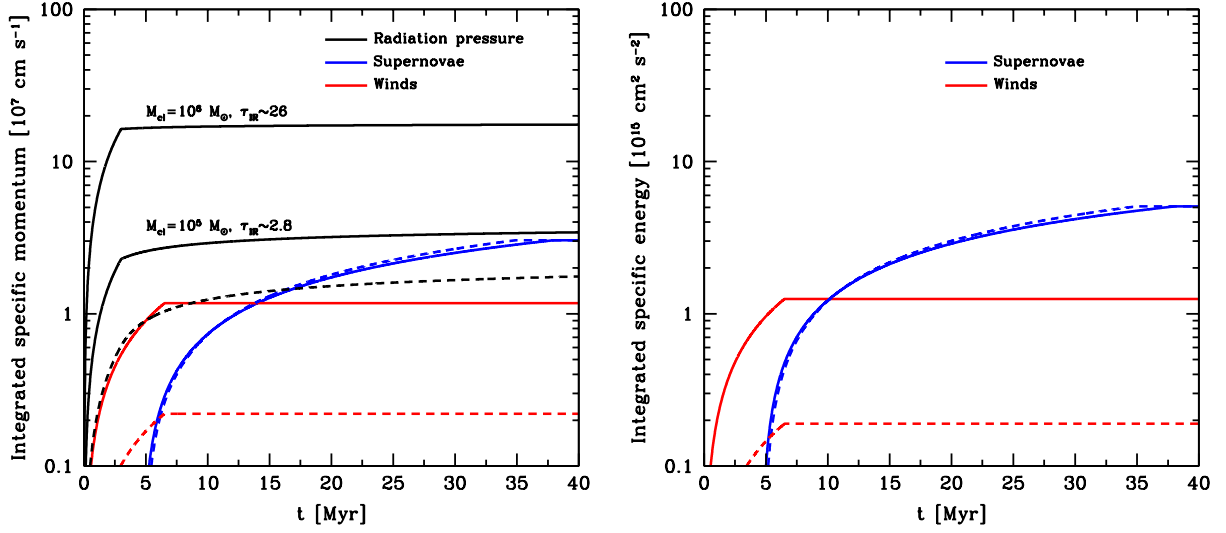


Figure 19.2: Integrated specific momentum (left) and energy (right) at $Z = Z_{\odot}$ (solid) and $Z = 0.01Z_{\odot}$ (dashed). This is the momentum and energy available as a function of time from a stellar population formed at $t = 0$ with unit mass. The momentum available in the radiation depends strongly on the degree of coupling between the radiation and the gas, as indicated by the τ_R parameter. From Agertz et al. (2013).

the limit where the cooling time is short, we can expect “momentum-driven feedback”. The other limit, in which $t_{\text{cool}} > t_{\text{dyn}}$ is known as “energy-driven feedback”.

We can understand the significant differences between these two classes by considering the energy available (see also Fig. 19.2). Imagine a cluster of stars has formed at the center of a cloud of gas. Let the momentum and energy injection rate from the stars be $\dot{p}_w$ and $\dot{E}_w$. The wind emanating from the stars sweeps up the surrounding gas into an outwardly moving shell. In the limit where all of the energy is radiated away, then the momentum of the shell is

$$p_{\text{sh}} = M_{\text{sh}} v_{\text{sh}} = \dot{p}_w t \quad (19.1)$$

where the last equality comes from momentum conservation. The kinetic energy of the shell in this case is

$$E_p = \frac{p_{\text{sh}}^2}{2M_{\text{sh}}} = \frac{1}{2} v_{\text{sh}} \dot{p}_w t. \quad (19.2)$$

In contrast, if none of the energy is radiated away, then the energy in the shell is simply $E_E = \dot{E}_w t$. The ratio between these two limits is:

$$\frac{E_E}{E_p} = \frac{2\dot{E}_w}{v_{\text{sh}} \dot{p}_w} \quad (19.3)$$

In the case of a baryonic wind (i.e., a wind comprised of baryons) we have $\dot{E}_w/\dot{p}_w = \frac{1}{2}v_w$ and so $E_E/E_p = v_w/v_{\text{sh}}$. In the case of a photon wind (i.e., a wind comprised of photons) $E = pc$ so that $\dot{E}_w/\dot{p}_w = c$ and $E_E/E_p = 2c/v_{\text{sh}}$. These are not small factors. For massive stars the baryon wind has velocities of several thousand km s^{-1} , while the expanding shell will typically have velocities of tens of km s^{-1} . So whether one is in the energy-driven or momentum-driven regimes has a very large impact on the energy available. And if one can tap into the energy available in the photon wind, then the energy available is enormous. Unfortunately, a photon wind attempting to “push” a cold baryonic shell of matter is generally Rayleigh-Taylor unstable. Simulating this process in detail is an active area of research.

We are now going to switch from considering feedback on small (cloud) scales to feedback on a galaxy-wide scale. In this case when we refer to a “wind” we refer to mass that is capable of leaving the entire galaxy. Let’s now consider what these limits imply for the expected behavior of the **mass-loading factor**:

$$\eta \equiv \dot{m}_w/\dot{m}_* \quad (19.4)$$

i.e., the ratio of mass-loss rate from the wind to the SFR. Note that the mass-loading factor is in principle measurable in the data, although it is in practice very difficult to measure the wind outflow rate. We are going to make an assumption that the wind on galactic scales has a velocity of order the halo velocity: $v_w \sim v_h$. Obviously, wind velocities lower than v_h will not escape the galaxy and so will not be a “wind” at all. And given that it is generally difficult to accelerate material to very high velocities on ~ 100 kpc scales, we can imagine that any wind that does develop will be just over the threshold, i.e., comparable to v_h .

With these assumptions we can estimate in the energy-driven wind limit the energy in the wind as

$$\dot{E}_w = \frac{1}{2}\dot{m}_w v_h^2 \quad (19.5)$$

while the energy injection rate from SNe is

$$\dot{E}_{\text{SN}} = 0.01\dot{m}_* f \epsilon_{\text{SN}} \quad (19.6)$$

where f is the coupling factor (a catch-all term that contains all the complexities of how the energy actually couples to the gas). We can deduce that in this limit:

$$\eta_E \propto v_h^{-2}. \quad (19.7)$$

In the case of a momentum-driven wind we have

$$\dot{p}_w = \dot{m}_w v_h \propto \dot{m}_* \quad (19.8)$$

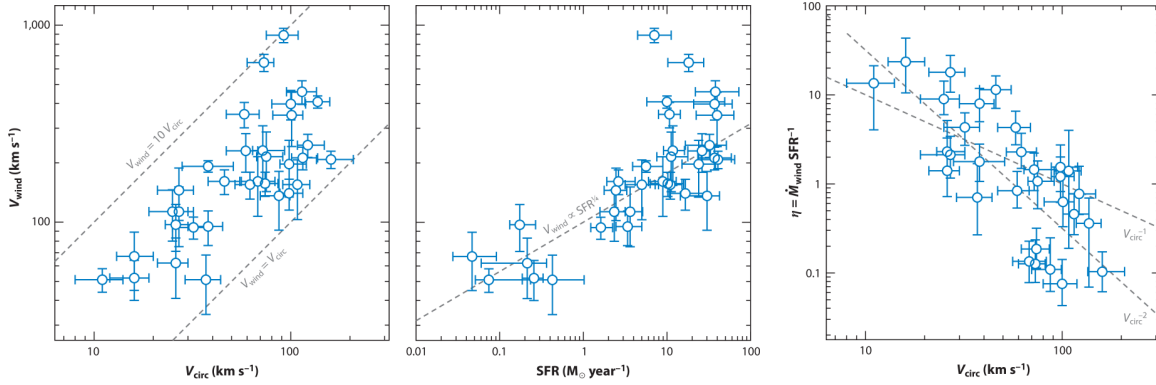


Figure 19.3: Velocity of outflowing cold gas vs. galaxy circular velocity (left panel) and SFR vs. galaxy circular velocity (middle panel). The right panel shows the mass loading factor vs. galaxy circular velocity, with the momentum-driven and energy-driven limits shown as dashed lines. From Xu et al. (2022).

so that

$$\eta_p \propto v_h^{-1}. \quad (19.9)$$

Many simulations model feedback-driven winds in a “sub-grid” manner in which the mass-loading factor is assumed to scale as one of the laws above. The approach is “sub-grid” because the actual physics of SNe explosions and the detailed radiation hydrodynamics occurs on too small of a scale for galaxy-wide simulations to resolve.

Observations of the outflow velocity and the mass loading factor for a sample of local starburst galaxies are shown in Fig. 19.3. The data more-or-less follow the basic scaling relations we outlined above.

We’re now going to consider the limit at which energy injection is sufficient to unbind the cold gas from a galaxy. We can expect this to occur on a dynamical time, i.e., $t_{\text{dyn}} \sim R/\sigma$ where R is the halo virial radius and σ is the average velocity dispersion of the halo.

Assuming the energy-driven limit, the condition for unbinding is:

$$E_b \approx \frac{GM_h M_g}{R} \approx \dot{E}_{\text{SN}} t_{\text{dyn}} \quad (19.10)$$

Let $M_g = f_g M_h$ where f_g is the cold gas fraction and M_h is the halo mass, and $\sigma^2 = 2GM_h/R$, so that $M_h^2 = \sigma^4 R^2 / 4G^2$. then

$$\dot{E}_{\text{SN}} \approx \frac{G f_g M_h^2 \sigma}{R} \quad (19.11)$$

$$\dot{E}_{\text{SN}} \approx \frac{f_g}{4G} \sigma^5 \approx 10^{-2} L_{\text{SB}} \quad (19.12)$$

where the last equality refers to the luminosity in the starburst and comes from the IMF and stellar population synthesis. The upshot is a scaling relation of the form:

$$L \propto \sigma^5 \quad (19.13)$$

this is the luminosity above which SNe are capable of unbinding the gas in the galaxy.

Nearly every galaxy contains a supermassive black hole (SMBH) at its center. When a BH accretes mass, it produces an associated accretion luminosity that is limited by the Eddington luminosity, which is the maximum luminosity at which radiation pressure forces just balance gravity (you will work through this concept in HW6). The upshot is that the Eddington luminosity is:

$$L_{\text{Edd}} = \frac{4\pi G m_{\text{BH}} m_p c}{\sigma_T} \quad (19.14)$$

where σ_T is the Thomson cross section and m_p is the proton mass. We can therefore assume that the energy provided by accretion onto a black hole is $\dot{E}_{\text{BH}} \propto L_{\text{Edd}} \propto m_{\text{BH}}$. Thus, energy-driven feedback from a black hole will produce a relation of the form:

$$m_{\text{BH}} \propto \sigma^5 \quad (19.15)$$

Momentum-driven winds yield a scaling of σ^4 for both the SMBH and SNe-based feedback mechanisms (setting $\dot{P} = GM^2/R^2$ and appealing to the Virial theorem to equate $M/R = \sigma^2$). These scaling dependencies are close to the observed relations.

The basic idea here is that σ is a reflection of the potential well depth. The galaxy (or SMBH) grows until it reaches the above scaling relation, at which point the associated feedback cuts off the fuel supply and limits further growth. The BH scaling relation is particularly remarkable and was initially unexpected, because it connects what is happening on AU scales (the SMBH) to what is happening on ~ 100 kpc scales (the halo potential). A recent version of the $M_{\text{BH}} - \sigma$ relation is shown in Fig. 19.4. The observed slope is in the range of 4 – 5 depending on the sample and analysis technique. It will probably be difficult to pin down the slope more precisely, but the basic point remains: some form of momentum or energy-driven feedback is likely important for shaping the observed relation.

A separate need for feedback arises when considering the state of galaxy clusters. The central galaxy in nearly all clusters is quiescent and generally comprised of very old stars. There is also very little if any cold ($T \lesssim 10^4$ K) gas. However, there is a huge reservoir of hot gas. This gas has cooling times in the central regions of $< 10^9$ yr, which means that – without additional input of energy – the hot halo would cool and collapse to form cold gas

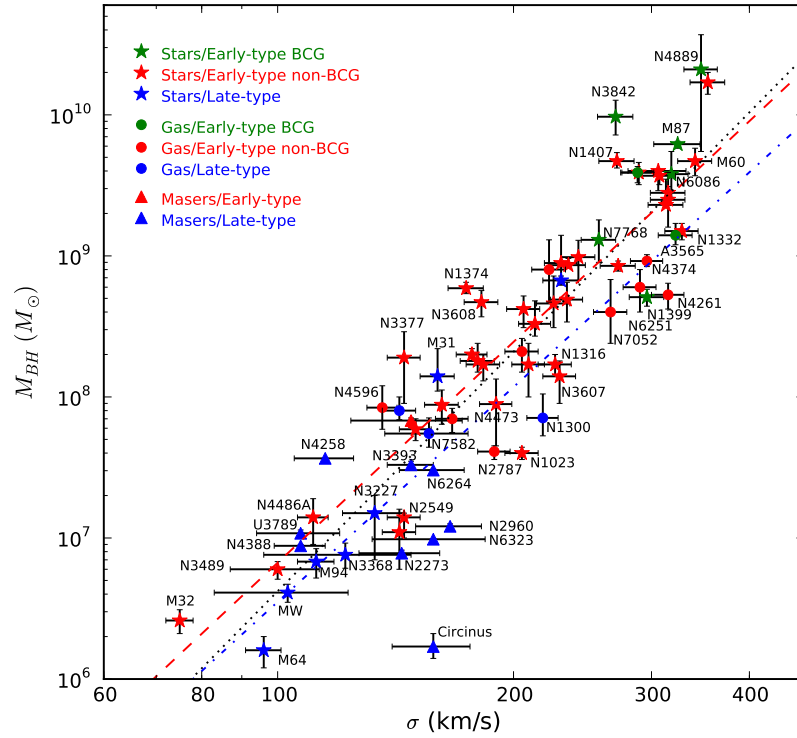


Figure 19.4: Relation between supermassive black hole mass and galaxy velocity dispersion for black hole masses estimated from stellar dynamics, gas dynamics, and maser emission. The dashed and dotted lines have a slope of ≈ 5 . From McConnell & Ma (2013).

and presumably stars. The fact that this is not observed is strong indication of a powerful source of feedback in clusters. While many ideas have been proposed, the current favorite is feedback from accretion onto SMBHs, though in a different form than discussed above. One idea is that SMBH accretion at low levels may generate enormous bubbles filled with cosmic rays which can efficiently heat the surrounding hot gas.

We have been speaking generally of feedback and stellar winds. It's time now to turn to a more detailed discussion of the available energy sources.

Let's first consider the energy available from massive stars. The basic sources are radiation pressure, ionizing radiation, stellar winds, and SNe. A significant advantage of the first three is that they are *early* in the sense that they provide energy as soon as the stars form, whereas the SN occurs only at the end of the star's life. These early forms of feedback are likely essential ingredients in keeping star formation inefficient on local (GMC) scales.

Consider radiation pressure. The momentum injection rate per unit mass is

$$\left\langle \frac{\dot{p}_{\text{rad}}}{M} \right\rangle = \frac{1}{c} \left\langle \frac{L}{M} \right\rangle = 20 \text{ km s}^{-1} \text{ Myr}^{-1} \quad (19.16)$$

where the last equality arises from considering the bolometric luminosity per unit stellar mass averaged over a stellar population. This can be summarized as follows: every gram of matter that goes into stars is capable – via radiation pressure – of accelerating another gram of matter to a velocity of 20 km s^{-1} over 1 Myr .

If we integrate over stellar lifetimes, this becomes:

$$\left\langle \frac{E_{\text{rad}}}{M} \right\rangle = 10^{51} \text{ erg } M_{\odot}^{-1} \quad (19.17)$$

That is a lot of available energy. Compare to SNe, which can provide 10^{51} erg per SN, but we only expect $\approx 1 \text{ SN}$ per $100 M_{\odot}$ of mass formed, so

$$\left\langle \frac{E_{\text{SNe}}}{M} \right\rangle = 10^{49} \text{ erg } M_{\odot}^{-1}. \quad (19.18)$$

The effect of ionizing radiation can also be significant. We take Q_H to represent the number of Hydrogen ionizing photons (with $h\nu > 13.6 \text{ eV}$) emitted by a population per second. At early times, before the most massive stars have evolved off the main sequence, we have:

$$\left\langle \frac{Q_H}{M} \right\rangle = 10^{47} \text{ s}^{-1} M_{\odot}^{-1} \quad (19.19)$$

It is not obvious from the number alone, but this is a lot of ionizing photons. These ionizing photons will heat the gas to $\approx 10^4 \text{ K}$ and create H II regions around the O stars. These regions will be very over-pressurized relative to their surroundings, and the net effect is to drive turbulence and eventually disrupt the cloud completely.

Finally, there is the hot stellar (baryonic) wind. For O stars, one observes $v_w \sim 10^3 \text{ km s}^{-1}$ and $\dot{m}_w \sim 10^{-6} M_{\odot} \text{ yr}^{-1}$. The momentum in these winds is relatively small compared to the radiation field, but the baryonic wind can still be energetically very important. One of the reasons for this was alluded to earlier – while the radiation has a lot of available momentum and energy, it can be difficult to efficiently *couple* that momentum and energy to the surrounding gas. If the radiation is able to escape, then it will have no impact on the surrounding medium.

There are also energy sources from lower-mass stars. Lower-mass stars are generally much less luminous and cooler, so they do not contribute meaningfully to the radiation pressure nor the ionizing radiation. However, they do have baryonic winds. The AGB stars in particular

have wind properties in the range $v_w \approx 10 - 30 \text{ km s}^{-1}$ and $\dot{m} \sim 10^{-5} - 10^{-4} M_\odot \text{ yr}^{-1}$. These winds are dense and cool, so they usually contain lots of dust and molecules. They can therefore cool very efficiently and so energy is not conserved. This means these winds are momentum-driven. These winds do not play an important role in terms of feedback. However, integrated over the lifetime of a stellar population, they are important insofar as they provide lots of mass back into the ISM that is then available for subsequent generations of star formation.

Feedback is a key ingredient for understanding the observed properties of both spiral and elliptical galaxies. But feedback is a complex process, involving accretion onto SMBHs, advanced phases of stellar evolution, and stellar explosions. Each of these is an active field of research with many unsolved questions, and so their detailed effect on the evolution of galaxies is similarly uncertain. Broadly, feedback is probably the single most important and uncertain aspect of our understanding of the formation and evolution of galaxies.

Chapter 20

Statistics of the Galaxy Distribution

Large surveys of galaxies enable analysis of the statistical properties of the population, including number counts and clustering. This has been a major area of study ever since the first galaxy redshift surveys in the 1970s. Statistics of the galaxy population has provided the primary measurements that enable inference of cosmological parameters and quantitative comparison to models and simulations. One of the most successful and important models for the clustering of galaxies is the halo model, which is a model for mapping galaxies to dark matter halos. In this chapter we will provide an overview of basic statistics used to quantify the large-scale distribution of galaxies as well as an overview of the halo model.

When we view the galaxy population in 3D (on-sky coordinates and redshifts) we observe a remarkable degree of structure on scales from at least 1 kpc to ~ 100 Mpc (recall that on larger scales the universe obeys the Cosmological Principle; i.e., it is homogeneous and isotropic). There have been many attempts to quantify this structure; by far the most popular are the n -point statistics.

The total number density of galaxies (e.g., above some luminosity or mass threshold) can be thought of as the 0-point statistic. For example, galaxies at the mass of the Milky Way have a space density of $n_{\text{MW}} \sim 10^{-2} \text{ Mpc}^{-3}$. The observable universe has a volume of 10^{12} Mpc^3 , so we can infer that there are approximately 10^{10} Milky Way-like galaxies in the universe.

The luminosity function (LF) is the corresponding 1-point statistic, and it quantifies the number density of objects per luminosity (or magnitude). See Fig. 20.1 for examples. The most common analytic model for the LF is the Schechter function:

$$\frac{dN}{dV dL} = \frac{dn}{dL} = \phi_* \left(\frac{L}{L_*} \right)^\alpha e^{-L/L_*} \quad (20.1)$$

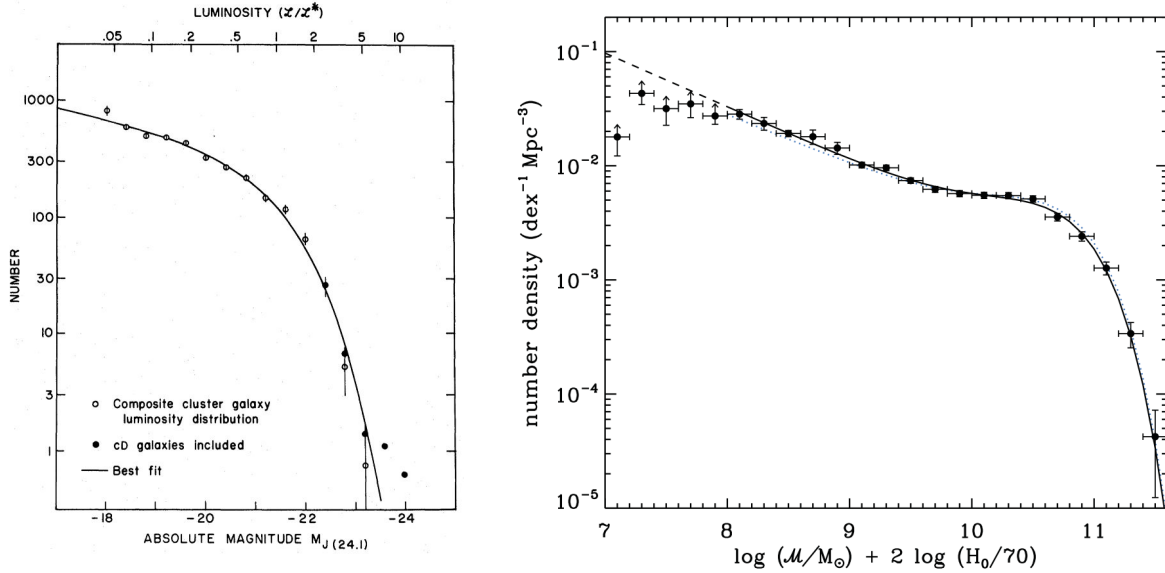


Figure 20.1: *Left panel:* Luminosity function of local galaxies, from Schechter (1976), in which the original “Schechter luminosity function” was first defined. *Right panel:* Stellar mass function of galaxies at $z \approx 0$, from Baldry et al. (2012). With a much broader range of masses, a more complex functional form is required to fit the data (two-part power-law and an exponential cutoff).

which is simply a power-law at $L < L_*$ and an exponential cutoff at $L > L_*$. The Schechter function has three key parameters: the normalization, ϕ_* , the characteristic luminosity, L_* , and the power-law slope, α . The data are well-fit by this simple function, although deviations have been noted at both the faint and bright ends. The deviation at the faint end is likely indicative of interesting physical processes, while the nature of the deviation at the bright end is still under debate (e.g., it could be an issue with measuring the extended envelopes of the most massive galaxies).

The LF varies depending on what band (wavelength) it is measured. At IR and X-ray wavelengths the LF appears to be closer to a broken power-law rather than a Schechter function.

The total number density of galaxies is an integral over the LF:

$$\int dn/dL dL \propto \phi_* \quad (20.2)$$

while the luminosity density is:

$$\int L dn/dL dL \quad (20.3)$$

and the average luminosity is:

$$\langle L \rangle = \frac{\int L dn/dL dL}{\int dn/dL dL} = f(\alpha) L_* \quad (20.4)$$

where $f(\alpha)$ is an order-unity function of the power-law slope; the key point is that the average luminosity is proportional to the characteristic luminosity, L_* .

In the optical one typically finds $\alpha \approx -1.2$, which implies that the luminosity density in the universe is finite (the luminosity density would be logarithmically divergent if $\alpha = -2$, although this is a somewhat pedantic point because the universe does not produce arbitrarily small galaxies). At $z = 0$ the characteristic luminosity in the B -band is $L_{B,*} \approx 10^{10} L_\odot$ which is approximately the luminosity of the MW. In absolute magnitudes this corresponds to $M_B^* \approx -20$.

A warning: the LF is often expressed in terms of magnitudes, i.e., dn/dM in which case the power-law index is flatter by one power (because $dn/dM \propto dn/d \ln L$). As always, be aware of factors of h . The normalization scales as h^3 , and so neglecting factors of $0.7^3 \approx 0.3$ can lead to significant errors and/or confusion.

How does one actually measure the LF? In principle this is trivial – count galaxies in bins of luminosity within a well-defined volume. In practice there are several issues that make this measurement difficult. There is large scale structure (LSS) in the universe, which means that the density of galaxies is not the same everywhere. This implies that one must survey sufficiently large volumes to overcome this extra variance. The variance due to LSS is often referred to in the literature as “cosmic variance” or “sample variance” and can be a serious source of uncertainty, especially for high redshift samples where the volumes probed are small (because of the longer exposure times required to study the faint high redshift universe).

A second issue is that all surveys of galaxies are flux-limited, which means that we can only observe galaxies above a given flux limit set by the telescope, exposure time, sky background, etc. This is one example of Malmquist bias and means that the effective volume sampled depends on the intrinsic luminosity of the galaxy (we will observe intrinsically more luminous galaxies to greater distances at a given flux limit). The classic solution to this problem is to weight each galaxy in the sample by $1/V_{\max}$, where $V_{\max}$ is the maximum volume over which that particular galaxy *could* have been detected, given the flux limit. This effectively up-weights the low-luminosity galaxies, and corrects for the flux limit effect. More modern approaches involve “forward-modeling” the population with Bayesian statistics.

Another issue at the bright end is that measurement uncertainty can have a very large effect on the shape of the exponential cutoff. This is an example of Eddington bias and can result

in a broadening of the LF at the bright end.

The Schechter parameters are often highly covariant, especially when data do not extend very far below L_* , which unfortunately is often the case. When measuring the LF at high redshift, it is common practice to fix the faint-end slope to a value estimated at lower redshift. Obviously this will introduce systematic uncertainties.

All of the machinery described above has also been applied to stellar mass functions, with the added step of first converting the light to mass via stellar population synthesis (as described in Chapter 17).

As we have seen in previous chapters, many galaxy properties are bimodal, with the quiescent, elliptical galaxies occupying one locus and the star-forming, spiral galaxies occupying the other. It is therefore common to measure the LF separately for these two populations, and to then measure evolution in the LF over time. These measurements provide unique insight into the evolution of the galaxy population(s) over cosmic time. In reality it might be better to imagine the full galaxy population as occupying a 2D surface of luminosity (or mass) and color (or SFR). Many erroneous results have resulted from selecting on one variable without taking proper account of the other. To give one example: a flux-limited survey is flux-limited in a particular band. Since quiescent and star-forming galaxies have different colors, this in general means that the effective mass-limit will be different for these two populations. If the depth is defined e.g., in the B -band, then one will observe star-forming galaxies to lower stellar masses than the quiescent galaxies.

Passive evolution has a large effect on the evolution of the LF. For example, the evolution of the LF from $z = 1$ to $z = 0$ for the quiescent population implies a dimming of the characteristic absolute magnitude, M^* by ≈ 1 mag. This is a direct consequence of the stellar evolution effects discussed earlier in this book.

Clustering statistics quantify the degree to which a set of points deviate from a random (Poisson) distribution. There are many options for how to quantify this. One of the simplest is the 2-point correlation function, ξ (also known as the autocorrelation function). Fig. 20.2 gives a sense of the structure we are attempting to quantify, and of how closely the galaxies trace the underlying dark matter.

We can define ξ as follows. Consider a universe in which the mean number density of galaxies is n . Then, for a statistically homogeneous universe, the probability of finding a galaxy within a volume element dV is $dP = ndV$. We define the correlation function, ξ , as the probability

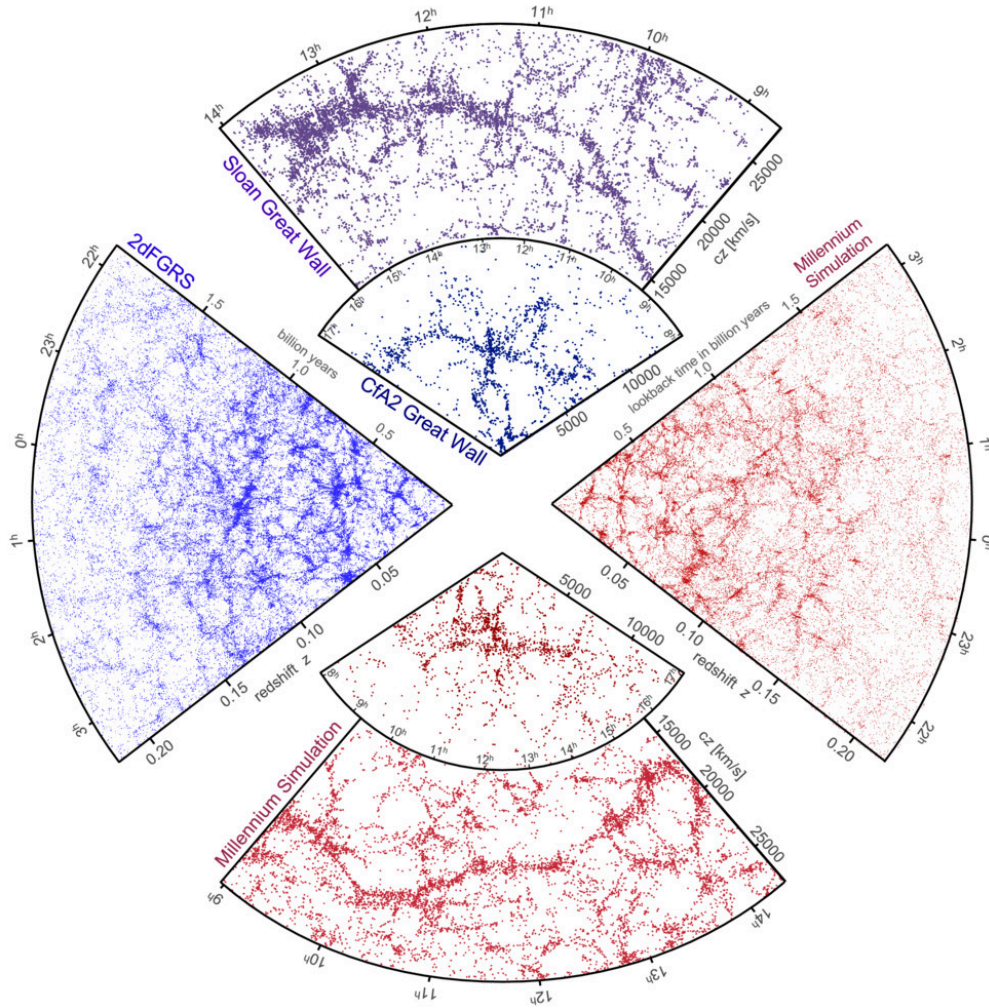


Figure 20.2: Large-scale distribution of galaxies and dark matter halos. The broad agreement between the spatial distribution of galaxies and dark matter is a key feature of the cold dark matter paradigm. Courtesy of the Virgo Consortium.

in excess of random (i.e., in excess of what would be expected in a statistically homogeneous universe) of finding a galaxy at a distance r from another galaxy:

$$dP = n[1 + \xi(r)] dV \quad (20.5)$$

We have seen the correlation function before, in our discussion of initial density fluctuations and the power spectrum, $P(k)$. Recall that the power spectrum and correlation function are Fourier transform pairs. The correlation function is a function of scale, $\xi(r)$, where one considers the volume element $dV = 4\pi r^2 dr$, and one measures the excess probability via pair counts at separation r .

In the limit of a Gaussian random field, the 2–point correlation function contains all of the

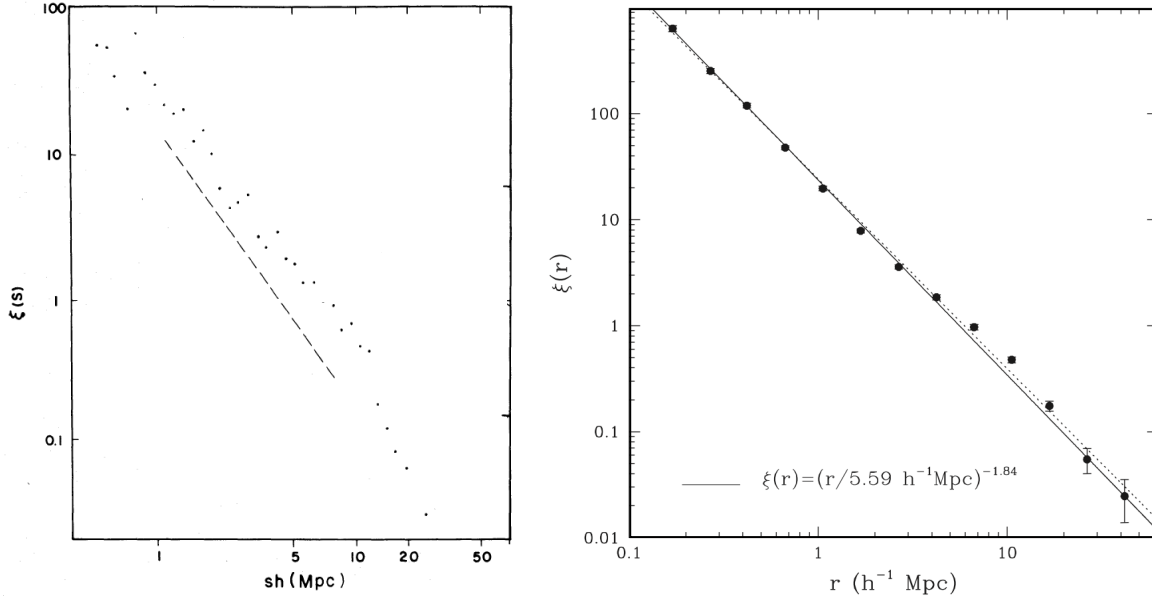


Figure 20.3: 2-point correlation function of galaxies, as measured in 1983 (left panel; Davis & Peebles 1983) and in 2005 (right panel; Zehavi et al. 2005). Remarkably, the dashed line in the left panel has a slope of -1.8 , essentially identical to modern measurements. Notice the small but statistically-significant departure from a power-law in the right panel at ≈ 10 Mpc (cf. Fig. 20.5).

information of the density field. However, the action of gravity produces a highly non-linear density field that results in useful information encoded in the higher-order moments, so one can generalize the above to the 3-point, 4-point, etc. The measurements become much more challenging, so in practice one does not generally go much beyond the 3-point.

One can also measure cross-correlation functions, in which one computes pairs from two different samples, e.g., red and blue galaxies, faint and bright galaxies, etc.

The measured correlation function is fairly well fit by a single power-law:

$$\xi(r) = \left(\frac{r}{r_0} \right)^{-\gamma} \quad (20.6)$$

with $r_0 \approx 4 - 5$ Mpc and $\gamma = 1.8$ for MW-like galaxies. The power-law is a reasonably good approximation from ≈ 10 kpc to ≈ 100 Mpc. Deviations from the power-law, though slight, have proven enormously consequential, as we will see below. In detail, the scale (r_0) and power-law index (γ) depend on everything – the galaxy luminosity, color, redshift, etc. See Fig. 20.3.

The underlying dark matter field is also clustered, with $r_0 \approx 4 - 5$ Mpc, although ξ_{DM} is not particularly well described by a single power-law. The dark matter halos are also clustered,

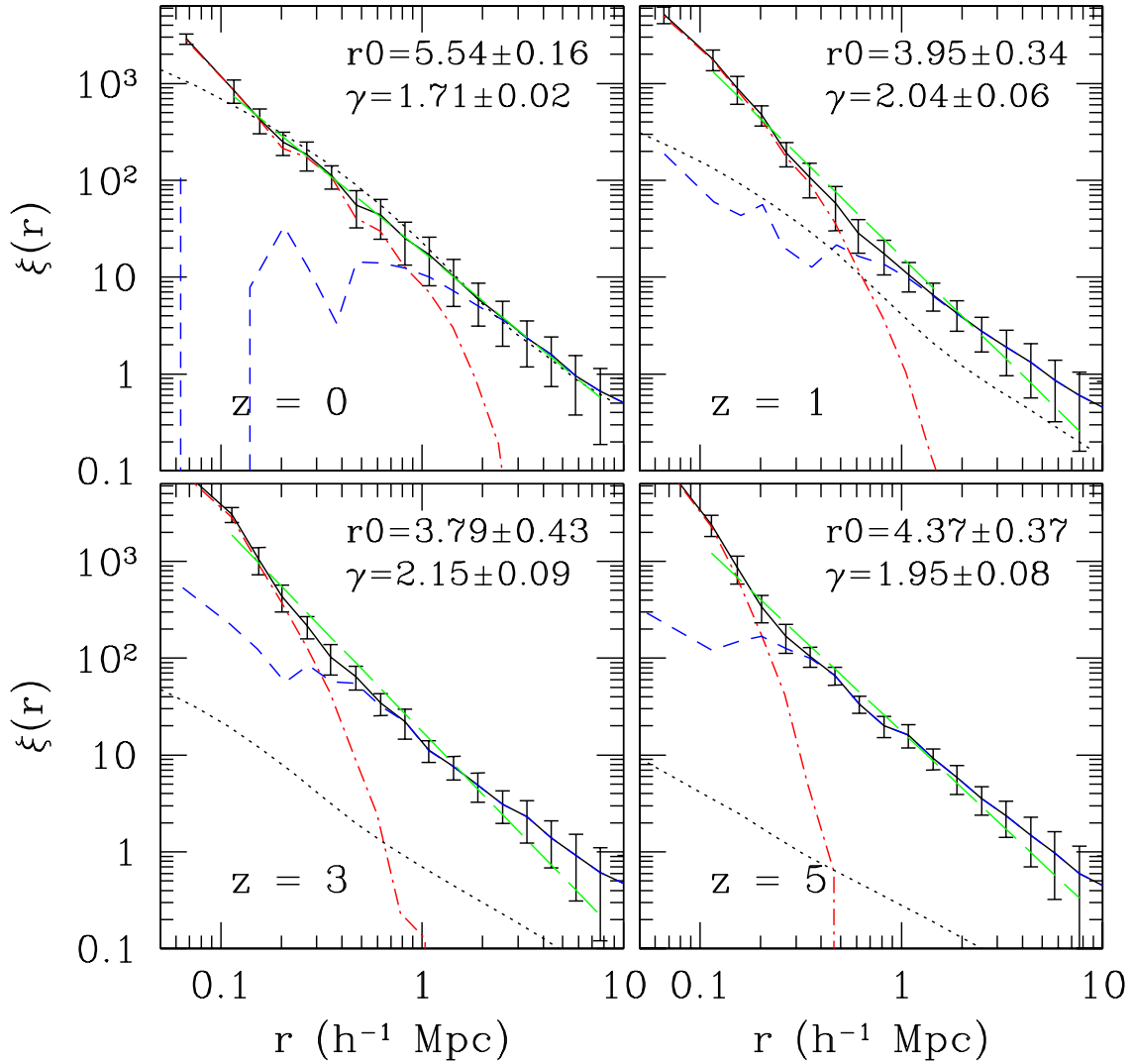


Figure 20.4: Correlation function for halos (solid lines with error bars) and dark matter (dotted lines) as a function of redshift. The 1-halo and 2-halo terms are shown separately as red and blue lines. A power-law fit is shown as a green line. Notice the much stronger evolution in the clustering of dark matter compared to the halos. The halo clustering deviates more strongly from a single power-law at higher redshifts. From Kravtsov et al. (2004).

and — critically — they do display power-law clustering. The clustering of the halos depends on halo mass and redshift. See Fig. 20.4.

There is an important concept known as bias, which is defined as:

$$b_X^2(r) \equiv \frac{\xi_X(r)}{\xi_{\text{DM}}(r)} \quad (20.7)$$

for any population X . One has for example the galaxy bias b_g and the halo bias b_h . The bias measures the clustering relative to the underlying dark matter density field. The bias is in principle a function of scale, although if it turns out that b does not depend on scale, then we refer to scale-independent bias. In the halos, bias increases with increasing mass, and for the galaxies bias increases with increasing luminosity, especially for $L > L_*$.

The fact that the correlation function of the overall dark matter density field is not a power-law while that of the dark matter halos is a power-law, is strong motivation for the **halo model**. The basic idea is that galaxies occupy halos with some occupation probability, and the clustering of galaxies is then a simple weighting of the halo clustering. Said another way: galaxy formation determines how galaxies occupy halos, and the large-scale structure of halos then determines the distribution of galaxies. See Fig. 20.5.

An example of a halo occupation distribution (HOD) is provided in Fig. 20.5. One can imagine this HOD as applying to a set of galaxies above some luminosity threshold. The halo mass below which $\langle N \rangle \approx 1$ is where one expects to find no halos hosting galaxies above the threshold. Above that threshold mass, galaxies either occupy the center of the parent halo (red dotted line) or orbit as satellites within it (pink dashed line). The total occupation is given by the sum of these two components (solid blue line). The HOD will therefore determine the resulting number density and clustering of galaxies, so one can imagine determining the HOD of a population given those observables. This is an active subfield in astronomy.

One can see how the galaxy bias is a function of the halo bias and the HOD from the following expression:

$$b_g = \frac{\int dM_h \, dn/dM_h \, b(M_h) \, N_g(M_h)}{\int dM_h \, dn/dM_h \, N_g(M_h)} \quad (20.8)$$

where dn/dM_h is the halo mass function, $b(M_h)$ is the halo mass-dependent bias, and $N_g(M_h)$ is the HOD. The equation above is just a number-weighted average of the halo bias.

The halo model offers a natural decomposition of the correlation function into two terms: a “1-halo” and a “2-halo” term, in which the galaxy pairs are either within the same halo or from two distinct halos. A clear and strong prediction of this model is that the apparent power-law behavior of the observed correlation function is only an approximation (and

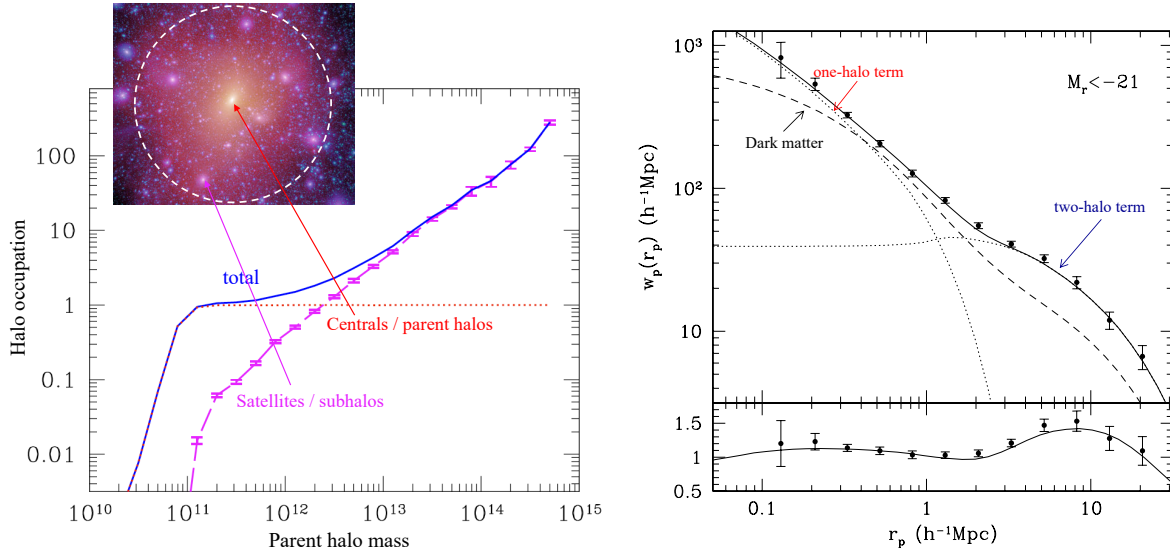


Figure 20.5: An overview of the halo model. *Left panel:* Example halo occupation distribution. The halo occupation is the mean number of objects as a function of halo mass, here shown separately for central and satellite galaxies. This occupation distribution is meant to represent a population of galaxies above a given luminosity threshold. From Kravtsov et al. (2004). *Right panel:* Projected correlation function measured in SDSS data (points), compared to a halo model fit (solid line). The fit is decomposed into 1-halo and 2-halo terms (dotted lines). The dashed line is the underlying dark matter correlation function. Notice the subtle, but statistically significant deviation from a power-law at 1 – 3 Mpc scales (the bottom panel shows the result divided by a best-fit power-law). From Zehavi et al. (2004).

something of a coincidence). In detail, one expects deviations from a power-law due to the transition between the 1-halo and 2-halo terms. This was a genuine prediction that was confirmed by subsequent measurements of clustering in the Sloan Digital Sky Survey (SDSS) Data. The right panel of Fig. 20.5 shows the measurements as well as the best-fit halo model and the separate 1-halo and 2-halo terms.

The halo model predicted that the deviation from a power-law would become more dramatic at higher redshift, and this too was confirmed by subsequent observations (see Conroy et al. 2006 for a summary). The basic physics will be familiar: the number of pairs within a halo depends on the balance between accretion of new galaxies and destruction by dynamical friction. It just so happens that at $z = 0$ this balance results in a 1-halo term that merges gracefully with the 2-halo term (which is set by the initial power spectrum plus quasi-linear structure growth). At higher redshift this balance shifts to favor a stronger 1-halo term.

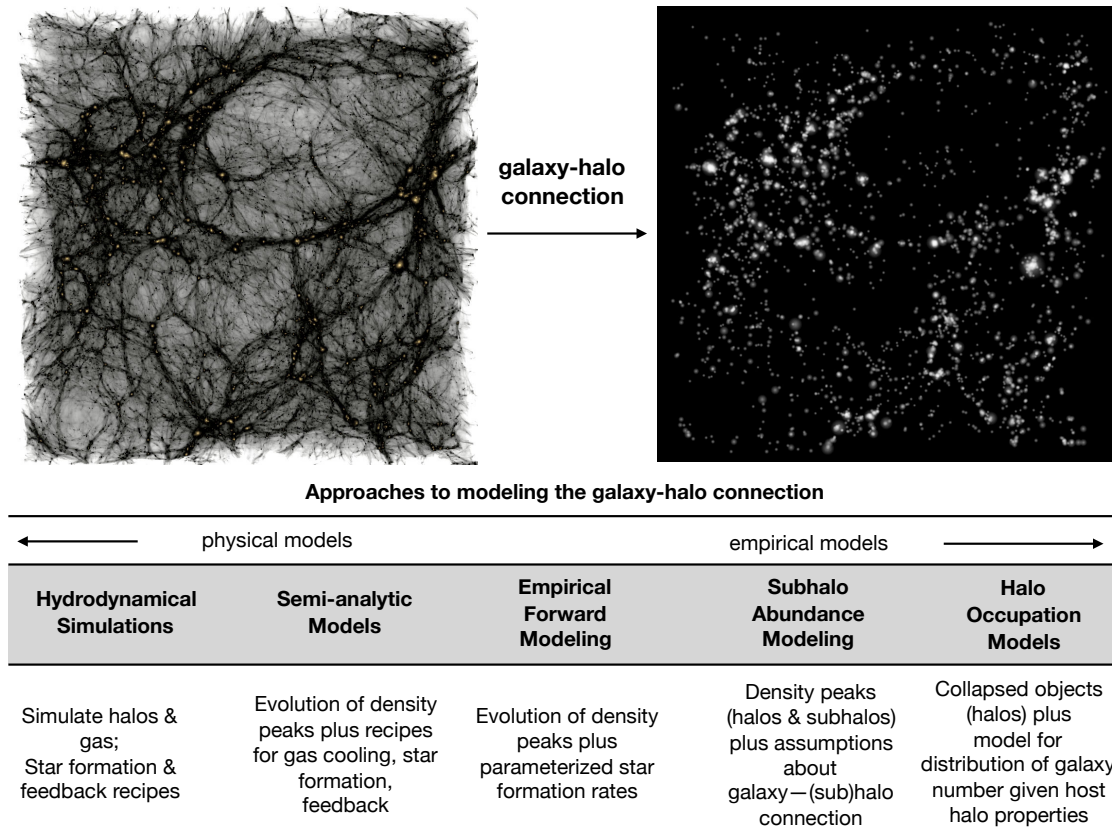


Figure 20.6: Overview of options for connecting galaxies to the cosmic web of dark matter halos. From Wechsler & Tinker (2018).

Moreover, because halos at a given mass are smaller at higher redshift, the 1-halo to 2-halo transition region, which reflects the characteristic halo size, shifts to smaller scales at higher redshift.

The halo model is only an approximation, and a relatively simple one. The key idea is that halo occupation depends solely on halo mass. Given this single assumption, it is remarkable that the model works so well at describing the large scale structure of galaxies. Extensions to the halo model have explored halo occupation distributions that depend on other halo parameters in addition to mass, including halo concentration, spin, large-scale environment, and mass accretion history. The halo model is one of several ways of connecting galaxies to halos; Fig. 20.6 provides an overview of the options.

Chapter 21

Successes and Challenges of Current Models

In this chapter we are going to put all of the ingredients we have assembled over the course into a coherent story of galaxy formation and evolution. This story will be explored further in your final project. We will also discuss several of the most significant open questions in the field. There are many excellent review articles on this subject. Several recent ones include Somerville & Dave (2014), Bullock & Boylan-Kolchin (2017), and Naab & Ostriker (2017).

Over the past 40 years a fairly detailed picture has emerged for the formation and growth of galaxies over cosmic time. A fundamental assumption is that the cosmological framework is well-understood and is comprised of cold dark matter (CDM) embedded in an expanding universe with a cosmological constant. The power spectrum of initial fluctuations is assumed to be a power-law, $P(k) \propto k$. With these simple (and well-tested, to varying degrees) assumptions one can, with the aid of computer simulations, predict the evolution of large scale structure (LSS) in the universe, from pc to Gpc scales. While the cosmological parameters, including the detailed shape of the power spectrum, are a continuing subject of interest for cosmologists, observations have reached a level of precision such that they are no longer a source of uncertainty for understanding how galaxies form and evolve. This was not the case 25 years ago. The upshot is that for many questions we can take the dark matter halos and merger trees as the backbone on which galaxy formation takes place.

As we have discussed in the latter half of this course, the emergence and evolution of galaxies requires understanding many interconnected processes including gas hydro and thermody-

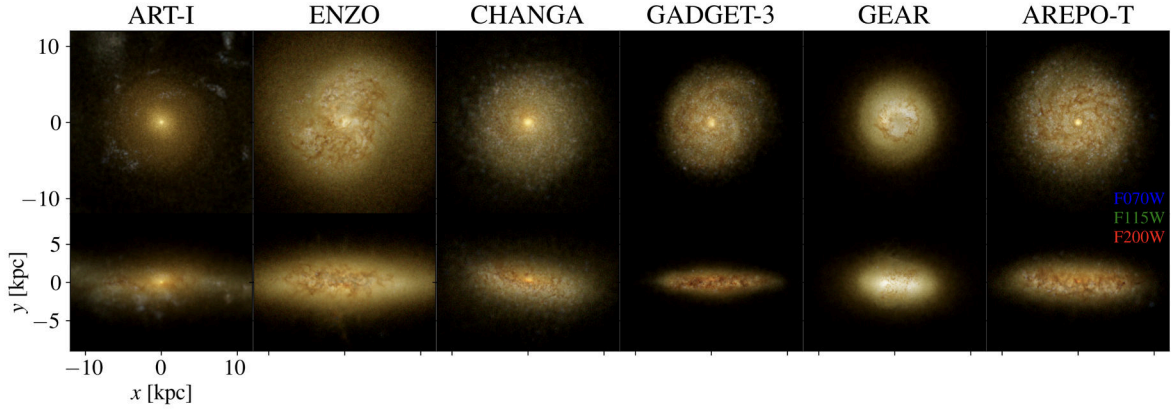


Figure 21.1: Simulated Milky Way-like galaxy at $z = 1$ as simulated by six popular galaxy simulation codes. The initial conditions at $z = 100$ are identical, so all differences can be attributed to numerical choices, assumptions, etc. From Jung et al. (2025).

namics, star formation and evolution, chemical evolution, and stellar and black hole feedback.

The gas hydro and thermodynamics are in principle straightforward problems — one must simply solve the fluid equations numerically and include the relevant heating and cooling functions. In practice, solving the fluid equations numerically is technically challenging for several reasons, not least of which is the enormous dynamic range involved. This means that one is always concerned that the solutions are not *numerically converged*, in other words, whether or not the results depend on the adopted resolution (or discretization). In detail, there are three techniques to discretize the fluid equations: mesh-based, smoothed particle hydrodynamics (SPH), and moving-mesh techniques. The first two can be thought of as solving the fluid equations in Eulerian and Lagrangian frames. The latter, which was an innovation in the past decade, is a hybrid scheme in which the mesh advects along with the fluid. There exist many codes to solve these equations (e.g., AREPO, GADGET, GASOLINE, ENZO, RAMSES, ART, GEAR, GIZMO, CHANGA). The AGORA Project (Kim et al. 2014, 2016, Jung et al. 2025) is an attempt to perform direct code comparisons to understand how robust galaxy formation simulations are (see Fig. 21.1).

To get a sense of how this field has evolved, consider that one of the first galaxy formation simulations was performed by Katz et al. (1996). This simulation had a spatial resolution of $5 - 10$ kpc and a mass resolution of $10^8 - 10^9 M_\odot$. Today, state-of-the-art simulations have a spatial resolution of $10 - 100$ pc and a mass resolution of $10^3 M_\odot$. Obviously this is a tremendous improvement, but the best simulations are still not resolving the structure of GMCs, supernovae blast waves, nor black hole accretion dynamics, nor are they simulating individual stars. See Fig. 21.2 for an example.

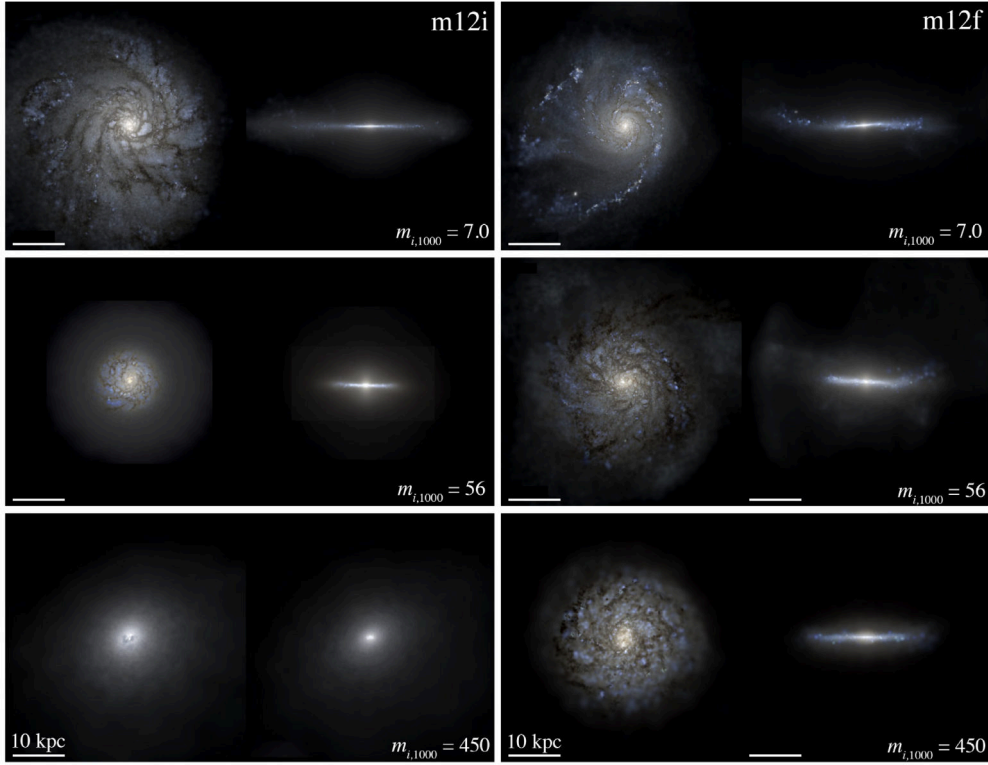


Figure 21.2: Dependence on the final properties of two Milky Way-like simulated galaxies as a function of simulation resolution limit. The mass resolution in the lower right of each panel is quoted in units of $1000M_{\odot}$. From Hopkins et al. (2018).

There are various tradeoffs one can make between simulating a large cosmological volume (in order to have a representative sample of simulated galaxies) and “zoom-in” simulations in which one focuses on the formation of a single galaxy. At fixed CPU time, the former will have lower resolution than the latter.

Owing to these shortcomings in resolution, the solution has been to resort to *sub-grid models* for physics occurring below the resolution limit, which includes, at a minimum, star formation and stellar and black hole feedback. Much effort has been devoted to developing realistic schemes for these important processes. More on this below.

Another approach to the problem of galaxy formation is to make the entire problem “sub-grid”. This approach is known as semi-analytic galaxy formation (SAM). The idea here is to use the dark matter merger trees and treat all baryonic processes, including cooling, heating, star formation, feedback, as well as the size evolution and morphological transitions of galaxies, as sub-grid. Modern models have 20–30 free parameters. They are not generally very predictive, but they are very fast, so one can run many realizations. They can be a

useful way to build intuition for the complex inter-dependencies of the problem. You will build a simple SAM for your final project.

We will now sketch out the basic story for galaxy formation. As dark matter halos collapse, gas is also being accreted, and will shock-heat to the halo virial temperature (although not always — there is “hot-mode” and “cold-mode” accretion of gas depending on the ratio of cooling and dynamical times in the halo). The shock-heated gas then has a cooling time that is a function of radius:

$$t_{\text{cool}} = \frac{3\rho(r)kT(r)}{2\mu m_H} \frac{1}{n_e^2(r)\Lambda[T(r), Z(r), \dots]} \quad (21.1)$$

As the gas cools it loses pressure support and collapses, often into a disk supported by angular momentum. The gas in the central regions will tend to cool first owing to the higher densities there.

Once the gas collects into a disk, it will continue to cool, and, as we learned in an earlier chapter, will settle into a multi-phase ISM with gas at 10^4 and 10^2 K. The coldest gas is susceptible to Jeans collapse and eventually stars will form when the densities become very large. The internal dynamics and collapse of GMCs cannot be captured in current galaxy-scale simulations, and so one must have a recipe for how stars form. There are several common schemes. One idea is to adopt a version of the empirical Kennicutt-Schmidt relation, $\dot{\rho}_* \propto \rho_g^{1.5}$. Another is to let stars form according to: $\dot{m}_* = \alpha_{\text{SF}} m_g / \tau_{\text{dyn}}$ where α_{SF} is a free parameter. As we’ve discussed before, what one chooses for the gas mass in the above expressions is important. Generally one makes available for star formation only the coldest and densest gas that one can resolve, ideally $T < 10^2$ K and $n > 10^3 \text{ cm}^{-3}$.

As we saw in an earlier chapter, feedback is critical to reproducing the observed scaling relations of galaxies as well as producing generally well-behaved galaxies. However, modeling stellar feedback is even more challenging than modeling star formation, and is the subject of many different sub-grid models. One example will give a sense of the challenge: if one deposits 10^{51} erg of thermal energy into a star forming region (to mimic the energy deposited by core-collapse SNe) then what usually happens in the simulation is the energy is immediately radiated away and so has little to no effect on the gas. This is a consequence of not being able to resolve the SNe blast wave properly. What to do? Some groups literally turn off gas cooling near the SNe for a time. Other groups numerically “kick” gas particles away from the star forming region in an attempt to smooth out the SNe effect. Obviously these are not great solutions. And how one implements these details can have a significant effect on the outcome. This issue has gotten somewhat better in recent years as the resolution has

increased to allow some modeling of the early SNe phases.

In the early 2000s the best models dramatically failed to reproduce the properties of the most massive galaxies (both their lack of star formation and their number densities). The problem was traced to “overcooling” at the cluster scale (left on its own, the hot cluster gas will cool in a fraction of a Hubble time and produce very high SFRs). In 2006, several groups began experimenting with essentially ad hoc and highly tuned models in which energy liberated from an accreting supermassive black hole (SMBH) coupled to (and heated) the surrounding gas on larger scales. At the simplest level, one relates the SMBH growth rate to the energy injection rate via: $\dot{e} = \eta \dot{m}_{\text{BH}} c^2$ where η is a free parameter. Developing and testing increasingly sophisticated sub-grid SMBH accretion models is a very active area of research.

On the bright side, the merger rates of galaxies, being controlled by the dark matter halo merger trees and dynamical friction, are a relatively stable and robust prediction of the models. Unfortunately, the merger rate is very difficult to measure observationally. Options include close-pair counts and using morphological disturbances as a proxy for mergers. Both require extensive testing and calibration with simulated data.

The above was a basic sketch of the basic ingredients necessary to make most galaxies. Now we will turn to some additional ingredients that may be important in certain regimes.

After recombination at $z \approx 1300$ the gas in the universe was neutral. This remained the case until $z \sim 7$, at which point the universe underwent its last great phase transition known as **reionization**. Much remains unknown about the origin of this transition, but it is believed to have occurred after the first galaxies emerged and began producing large quantities of ionizing radiation (see Fig. 21.3). Reionization causes the temperature of the intergalactic medium (IGM) to increase to $\approx 10^4$ K. One important implication of this increase in temperature is that halos with $T_{\text{vir}} < 10^4$ K will not be able to accrete gas from the IGM, and the gas within those halos may become unbound if that gas is also heated to $> 10^4$ K. At $z = 7$ this corresponds to a halo mass of $4 \times 10^9 M_{\odot}$. Such halos evolve into halos of mass $3 \times 10^{10} M_{\odot}$ by $z = 0$. The implication is that halos below this mass threshold should not only be devoid of gas but will also likely be devoid of stars formed at $z < 7$. “Reionization quenching” may be the origin of the ancient stellar populations in the ultra-faint dwarf galaxies.

Our story has neglected a number of potentially important physical effects including magnetic fields, cosmic rays, thermal conduction, radiation transport, and other heating sources

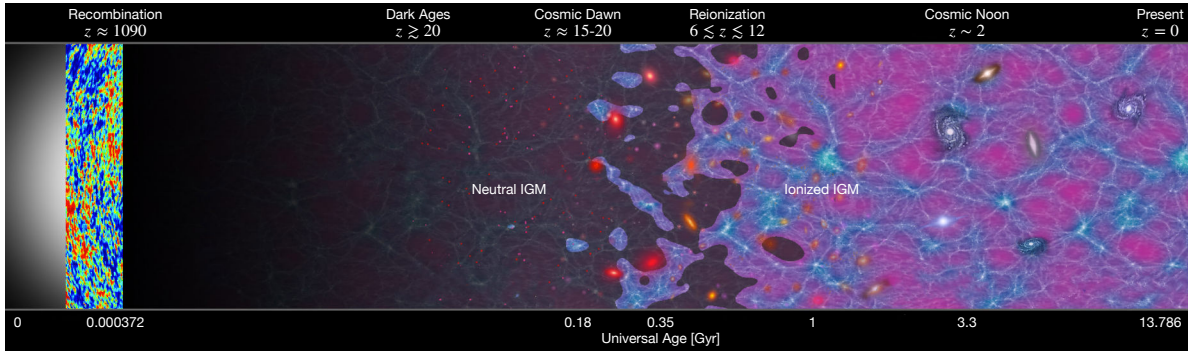


Figure 21.3: Sketch of the history of the universe with a focus on the evolution of the state of the IGM. From Robertson (2022).

including AGB stars and low-mass X-ray binaries. Various groups have experimented with each of these effects in different contexts. In some cases they can have leading order effects, though for the basic story they are usually sub-dominant.

We said earlier that galaxy formation takes place on the backbone of the dark matter, but that is not the whole story. Galaxy formation can modify the distribution of dark matter within halos in important ways, either via the adiabatic contraction of the gas (which can increase the density of dark matter relative to its initial state), or via impulsive feedback (which can decrease the density of dark matter). The extent to which galaxies can sculpt the underlying dark matter distribution has enormous consequences for a class of “problems” with CDM predictions on small scales.

There are many related “small-scale problems”. A common thread is apparent tension between CDM simulation predictions and observations on the scale of small galaxies/halos (e.g., $M_h \lesssim 10^{10} M_\odot$). The first such problem was known as the “missing satellites problem”. The first generation of CDM simulations that reliably resolved the substructure within MW-size halos predicted a huge number of subhalos, whereas we only observe ~ 30 satellite galaxies in the MW halo. The accepted solution to this problem is that galaxy formation becomes increasingly inefficient in smaller halos, until a threshold is reached where halos have no luminous counterpart. Another problem on small scales is the core-cusp problem. CDM simulations predict a density profile in the inner regions of halos of $\rho \propto r^{-1}$, i.e., a “cusp”. Observations tend to favor a core: $\rho \propto r^0$. One solution to this problem was alluded to above: namely that galaxy formation has the capacity to alter the density slope in the inner regions via feedback (the basic physics is akin to violent relaxation). Whether or not this explanation is the answer is an outstanding question. See Fig. 21.4.

There are other problems on small scales, but they all have the same basic flavor. One

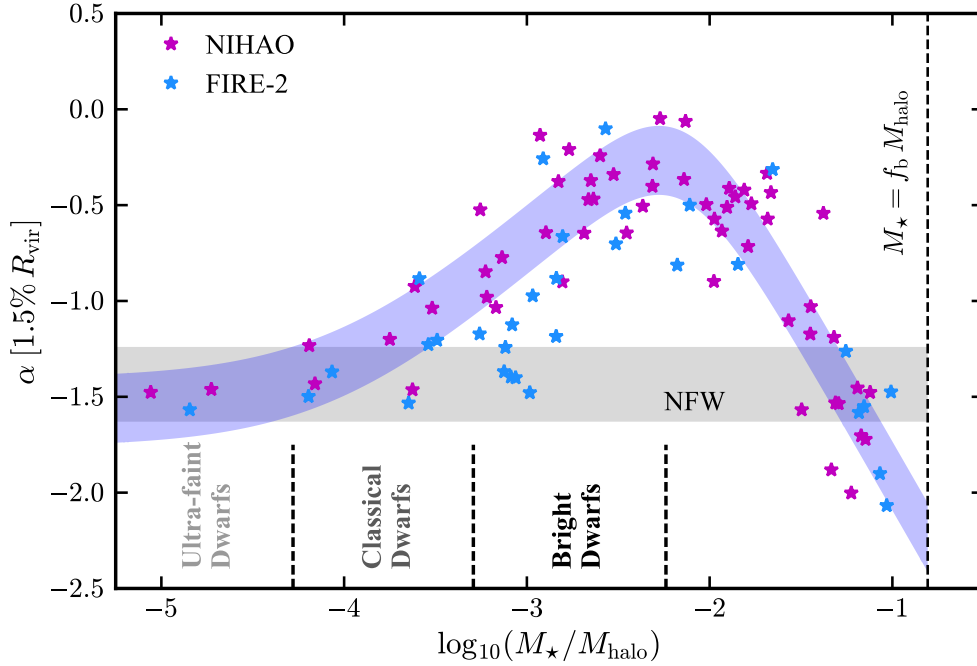


Figure 21.4: Slope of the dark matter density profile, measured at 1.5% of the virial radius, as a function of the ratio of stellar-to-halo mass. Note that this is the *local* slope at that radius, which for an NFW profile is ≈ -1.5 , steeper than the -1 asymptotic inner slope. Points and blue shaded region are results from two modern galaxy formation simulations. Profiles shallower than NFW are expected for bright dwarfs due to efficient feedback altering the initially steep profiles. At sufficiently low stellar masses, simulations predict that the density profiles should reflect the initial NFW profile. From Bullock & Boylan-Kolchin (2017).

reason these problems receive so much attention is that if galaxy formation cannot explain the tension, then perhaps we may learn something novel about the nature of dark matter. For example, warm dark matter predicts much less structure on small scales, and may in principle provide a suitable explanation to at least some of the small-scale problems. Because of the implications for cosmology and dark matter, this is a very active area of research.

Modern simulations of galaxy formation are now quite sophisticated and are generally able to reproduce many if not all of the key observables that they were designed to reproduce. However, the sub-grid models for star formation and stellar and black hole feedback are largely untested and contain many free parameters. A major area of ongoing work is to improve those sub-grid models, either by careful calibration against observations or via higher resolution simulations.

Historically, “star particles” in simulations were assumed to consist of a fully-sampled IMF of often thousands of stars. As the mass resolution in simulations pushes ever lower, eventually simulations will have to reckon with the truly discrete nature of stars. That should produce interesting phenomena especially in low-mass galaxies.

The circumgalactic and intergalactic media (CGM; IGM) have remained largely inaccessible observationally, and so models are not well-tested in this regime. This situation is likely to change in the coming decade, as a suite of new observatories, instruments, etc. seek to probe the CGM and IGM in novel ways.

Rubin/LSST, which starts survey operations soon, will deliver a complete census of the faintest galaxies orbiting the Milky Way, and will therefore provide important constraints on the limits of galaxy formation at the faint end.

JWST has delivered on its promise to detect and characterize galaxies at $z > 10$. Galaxies at these early epochs appear to be far more abundant than models predicted. At slightly lower redshifts ($4 < z < 8$) *JWST* is finding far more massive and evolved galaxies than expected. Both observations point toward much more rapid and efficient galaxy formation at early epochs than predicted. Implications for galaxy formation models are still being worked out.