

Системы и методы биржевой торговли Forex

Перри Кауфман

Глава 2. Базовые концепции и расчеты

Экономика — не точная наука: она состоит из одних законов вероятности. Поэтому самым благоразумным инвестором является тот, кто следует общему курсу, который «обычно» является правильным, и кто избегает действий и идей, которые «обычно» являются неправильными.

Л. Ангас

Новые технологии дают нам ощущение безопасности. Мы моментально можем получать данные из любых точек мира, у нас есть программы, которые мгновенно выполняют самые сложные расчеты, и мы в любое время можем связаться с кем угодно.

Как предсказывал Айзек Азимов, придет время, когда мы уже не будем знать, как делить в столбик, поскольку повсюду будут миниатюрные, управляемые голосом компьютеры. Может быть, мы даже разучимся складывать, и это будут делать за нас. Мы просто будем полагать, что ответ правилен, потому что компьютеры не ошибаются.

В некоторой степени это происходит уже сейчас. Не все проверяют расчеты электронных таблиц вручную, чтобы убедиться в их правильности, прежде чем двигаться дальше. И далеко не все распечатывают промежуточные результаты компьютерных вычислений, чтобы проверить их точность. Компьютеры не совершают ошибок, но люди их делают.

С появлением программного обеспечения и торговых платформ, сделавших анализ цен одновременно легче и сложнее, мы больше не думаем о том, как, собственно, работает скользящая средняя или линейная регрессия. Несколько лет назад мы рассматривали корреляцию между инвестициями только в случае крайней необходимости, потому что для этого требовались слишком сложные и длительные вычисления. При этом нельзя было знать наверняка, что вы не допустили ошибки, пока кто-то еще не проделает те же вычисления повторно. Теперь мы стоим перед другой проблемой: если все делает компьютер, мы перестаем понимать, чем скользящая средняя отличается от линейной регрессии. Не видя исходных данных, мы не замечаем ошибки, вызываемой аномальным отклонением, или того, что в цене акции не было учтено дробление. Не изучая каждую гипотетическую сделку, мы теряем возможность видеть, как проскальзывание может превратить прибыль в убыток.

Чтобы избежать пробела в знаниях, необходимых для создания прибыльных торговых стратегий, в этой главе объясняются базовые инструменты торговли. Те из вас, кто уже знаком с ними, могут пропустить главу, остальным же следует удостовериться, что они могут выполнять вычисления вручную, даже если используют электронную таблицу.

Полезное программное обеспечение

В Excel встроены многие необходимые функции, например стандартное отклонение, и их можно использовать в любое время. Более сложные статистические функции необходимо подключать дополнительно как надстройки, но они поставляются с Excel бесплатно. Это гистограммы, регрессионный анализ, *F*-критерий, *t*-критерий, *z*-критерий, анализ Фурье и различные методы сглаживания. Чтобы подключить эти надстройки в Excel 2010, перейдите в Файл/Параметры/Надстройки и выберите «Пакет анализа». Вам также потребуется надстройка «Поиск решения». После подключения, которое занимает всего несколько секунд, к этим функциям можно получить доступ в меню «Данные» в верхней части экрана. Опции «Анализ данных» и «Поиск решения» будут находиться в меню справа.

Есть и другие очень полезные и легкие в использовании статистические программы, различающиеся по сложности и по цене. Одним из самых выгодных приобретений может оказаться *Pro-Stat* от Poly Software (polysoftware.com). Примеры в данной главе построены с использованием как Excel, так и *Pro-Stat*.

О данных и усреднении

Закон больших чисел

Начнем с начала, с закона больших чисел — очень неправильно понимаемого и неверно цитируемого принципа. В торговле на закон больших чисел чаще всего ссылаются, когда ожидается, что ненормально продолжительная последовательность убытков будет компенсирована равным периодом прибыли. Одинаково неправильно ожидать, что рынок, который в настоящее время переоценен или перекуплен, станет затем недооцененным или перепроданным. Это совсем не то, что гласит закон больших чисел: в большой выборке большинство событий рассеивается близко к среднему значению таким образом, что типичные значения существенно превосходят нетипичные, делая их незначимыми.

Этот принцип проиллюстрирован на рис. 2.1, где число средних величин чрезвычайно велико, поэтому добавление небольшой неправильной группировки с одной стороны средней группы почти нормальных данных не нарушает равновесия. Это все равно, что влияние единственного пассажира на аэробус. Ваш вес не имеет значения для самолета и ни на что не влияет, даже когда вы перемещаетесь по салону. Длительная череда прибылей, убытков или необычно продолжительное движение цен в одном направлении — это просто редкое, аномальное событие, которое со временем сглаживается подавляюще большим количеством нормальных событий. Дальнейшее описание этой проблемы и ее влияния на торговлю приведено в главе 22 «Методы азартной игры — теория выбросов».

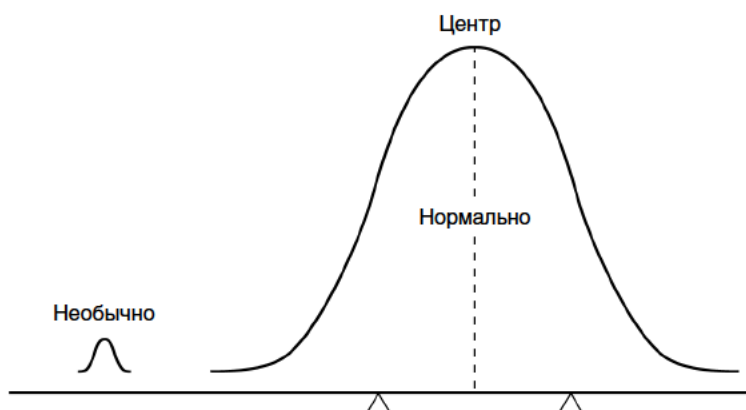


Рис. 2.1. Закон больших чисел. Нормальные события значительно превосходят необычные. Для создания равновесия нет необходимости в чередовании экстремальных событий — одно выше, другое ниже

Данные в пределах и за пределами выборки

Строгое соблюдение процедур тестирования требует разделения данных на множества *в пределах* и *за пределами* выборки. Подробнее это описано в главе 21 «Тестирование систем». Пока же рассмотрим самые важные моменты. Все тесты грешат чрезмерной подгонкой данных, однако нет иного способа узнать, работает ли идея или система, кроме как протестировать ее. Отбирая данные для последующего использования в проверке достоверности теста, вы повышаете шансы на то, что ваша идея будет работать, прежде чем проверять ее на реальных деньгах.

Существует множество способов отбора данных в пределах выборки. Например, если у вас есть история цен за 20 лет, вы можете использовать первые 10 лет для тестирования, а вторые 10 лет оставить для проверки достоверности теста. Конечно, рынки с течением времени изменяются — они становятся более волатильными и могут быть более или менее склонными к трендам. Поэтому лучше всего использовать перемежающиеся периоды данных в пределах и за пределами выборки, в виде двухлетних интервалов например, при условии, что вы не будете заглядывать в данные во время периодов вне выборки. Чередование этих периодов может создавать проблему в случае непрерывных долгосрочных трендов, но она решается в главе 21.

Самым важным фактором, который нужно учитывать при резервировании данных за пределами выборки, является то, что вы получаете только один шанс использовать их. После того как вы создадите правила для торговой программы, наступает время прогнать эту программу через неизвестные данные. Если результаты будут успешны, вы сможете торговать с помощью системы, но, если результаты плохой, вся ваша работа идет насмарку. Вы не можете изучать причины плохих результатов и изменять метод торговли, чтобы добиться лучшего. Это означало бы использование *обратной связи*, после чего данные за пределами выборки считаются недостоверными. Вторая попытка всегда лучше, но здесь уже наблюдается чрезмерная подгонка.

Сколько нужно данных

Статистики скажут: «Чем больше, тем лучше». Чем больше данных для тестирования, тем надежнее результаты. Техническому анализу повезло, что он основан на совершенном множестве данных. Каждая цена, зарегистрированная биржей, будь то IBM на закрытии в Нью-Йорке 5 мая или цены процентных ставок по евродолларам в 10:05 в Чикаго, является подтвержденной точной величиной.

Помните, что, когда вы используете для разработки данные в пределах и за пределами выборки, нужно как можно больше данных. У вас останется лишь половина комбинаций и моделей, когда вы отберете 50% данных.

Экономические данные

Большинство статистических данных не столь своевременно, не столь точно и не столь надежно, как цены и объемы акций, фьючерсов, ETF и других торгуемых на бирже продуктов. Экономические данные, например индекс промышленных цен (Producer Price Index — PPI) или число строящихся жилых домов, публикуются в виде среднемесячных значений (индикаторов) и могут корректироваться с учетом сезонных колебаний. Среднемесячное значение представляет собой широкий диапазон чисел. В случае PPI некоторые производители могли заплатить меньше среднего значения за предшествующий месяц, а некоторые больше, при том что среднее число (индикатор) составило, например, +0,02. Отсутствие диапазона значений, или стандартного отклонения составляющих величин, снижает полезность информации. Статистические данные часто пересматриваются в следующем месяце, иногда весьма существенно. При работе с еженедельными выпусками данных Министерства энергетики нужно знать точную историю опубликованных цифр, включая все корректировки, если вы хотите создать метод торговли, который реагирует на эти отчеты. Вы можете столкнуться с тем, что намного легче найти исправленные данные, но это совсем не то, что действительно нужно.

При использовании экономических данных нужно знать, когда они выпускаются. США в этом очень точны и оперативны, но другие страны могут опаздывать с выпуском данных на месяцы и годы. Если вашей программе требуется ввод ежемесячных данных из CRB Yearbook, не забудьте уточнить, когда эти данные стали фактически доступными.

Ошибка выборки

При построении экономического индикатора необходимо иметь достаточно много данных, чтобы сделать этот индикатор точным. Поскольку многие статистические данные являются выборочными, особое внимание уделяется аккумулярованию достаточного количества представительных данных. То же относится и к ценам. Усреднение нескольких цен или анализ небольших движений рынка демонстрирует менее достоверные результаты. Трудно нарисовать точную картину на основе очень небольшой выборки.

Когда используются небольшие, неполные или репрезентативные множества данных, можно найти ошибку приближения, или точность выборки с помощью

стандартного отклонения. Большое стандартное отклонение говорит о чрезвычайно рассеянном множестве точек, что, в свою очередь, делает среднее менее представительным в отношении данных. Этот процесс называют *тестированием значимости*. Точность увеличивается, когда количество данных становится больше, и величина *ошибки выборки* становится пропорционально меньше:

$$\text{Ошибка выборки} = \frac{1}{\sqrt{\text{Количество единиц данных в выборке}}} = \frac{1}{\sqrt{N}} \text{ или } \frac{1}{\sqrt{(N)}}.$$

Следовательно, при использовании только одной единицы данных ошибка выборки составляет 100%; при четырех единицах ошибка составляет 50%. Размер ошибки важен для надежности любой торговой системы. Если в системе проведено только четыре сделки, не важно, прибыльных или убыточных, очень трудно сделать сколько-нибудь надежные выводы о будущих результатах. Должно накопиться достаточное количество сделок, чтобы можно было с уверенностью говорить о небольшом коэффициенте ошибки. Чтобы уменьшить ошибку до 5%, нужно провести 400 сделок. Это представляет проблему для очень медленных методов следования за трендом, где может совершаться лишь две-три сделки в год. Чтобы компенсировать это, можно применить идентичный метод на нескольких рынках и использовать все проведенные на них сделки вместе (больше об этом в главе 21).

Репрезентативные данные

Количество данных — хороший показатель их полезности, при этом, однако, данные должны представлять как минимум один бычий рынок, один медвежий рынок и несколько периодов бокового движения. Если рынков каждого типа больше чем один, еще лучше. Если вы хотите использовать дневные значения индекса S&P за 10 лет с 1990 по 2000 г. или цены 10-летних казначейских нот за последние 25 лет до 2010 г., вы увидите только бычий рынок. Здесь торговая стратегия была бы прибыльной всякий раз на стороне покупателя, если держать позицию достаточно долго. Но, если не включить сюда множество других моделей поведения цены, у вас не получится стратегия, способная пережить падение рынка. Ваши результаты будут нереалистичными.

Данные, утратившие полезность

Бывают очевидные случаи, когда на рынке акций или фьючерсов происходит структурное изменение и текущие данные начинают отличаться от исторических. Эволюция General Electric из изготовителя лампочек в крупный финансовый институт является примером именно такого структурного изменения. Ее превращение обратно в производственную компанию, о котором объявлено в 2010 г., может стать другим структурным изменением. У компании, начавшей свою деятельность в США (McDonald's, например), но вышедшей далеко за национальные границы, структурное изменение также проявляется в поведении цены. На рынке иностранных валют мы наблюдали, как отдельные европейские валюты сначала были связаны валютным соглашением, а затем слились в единую валютную единицу евро.

Действительно ли важно включать исторические данные в тесты, когда эти данные представляют другой профиль деятельности компании или иную геополитическую ситуацию? В идеале ваша стратегия будет надежной тогда, когда сможет приспосабливаться к изменению профилей, постоянно демонстрируя положительную доходность в течение продолжительного периода тестирования. Статистики здесь правы — действительно, чем дольше, тем лучше. Ведь компании и рынки продолжают развиваться, и ваша программа должна будет постоянно приспосабливаться.

Если вы очень быстрый трейдер, то можете ограничить тестирование гораздо более короткими периодами. Если вы торгуете раз в день, за пять лет у вас наберется 1250 сделок, за 10 лет — 2500 сделок. Если ваша торговая стратегия прибыльна на протяжении 2500 сделок, то вы решили проблему ошибки выборки. Однако вы, возможно, не включили данные, которые являются репрезентативными для других рынков и множества других моделей поведения цены. Даже при большом количестве сделок нужны тесты, охватывающие много лет, чтобы продемонстрировать надежность системы.

Лучше перестраховаться

Важно помнить, что точность тестирования зависит и от количества используемых данных, и от числа сделок, проведенных системой. Если оценки убытков ненадежны, вы подвергаете свои инвестиции опасности.

О средних значениях

В работе с числами нередко бывает необходимо использовать репрезентативные значения. При решении задачи может подставляться диапазон значений или среднее, чтобы превратить отдельную цену в общую характеристику. Среднее (среднее арифметическое) многих значений может быть предпочтительным заместителем любого отдельно взятого значения. Например, средняя розничная цена одного фунта кофе на Северо-Востоке более значима при расчете стоимости жизни, чем цена в любом отдельно взятом магазине. Однако не все данные можно комбинировать или усреднять без потери смысла. Среднее всех цен, взятых за один день, ничего не скажет ни о каком отдельном рынке, который является частью среднего значения. Усредняя цены несвязанных вещей, например пачки кукурузных хлопьев, нормо-часа в авторемонтной мастерской и немецкого индекса DAX, мы получим число весьма сомнительной ценности. Среднее группы чисел должно иметь какой-то полезный смысл.

Среднее может запутать и по-другому. Рассмотрим цены на кофе, выросшие в течение года с \$ 0,40 до \$ 2,00 за фунт. Средняя цена этого продукта составляет \$ 1,20, однако она не учитывает время, в течение которого кофе продавался по разным ценам. В таблице 2.1 рост цены на кофе разделен на четыре равных интервала. Как видно, время, проведенное на каждом из этих уровней, обратно пропорционально росту цен. Иными словами, цены находились на более низких уровнях в течение более длительных, а на более высоких — в течение более коротких периодов, что совершенно нормально для поведения цены.

Таблица 2.1. Взвешивание среднего значения

Цены изменяются		Средняя в течение интервала	Всего дней в интервале	Взвешенное	1/a
от	до				
40	80	$a_1 = 60$	$d_1 = 100$	6000	0,01666
80	120	$a_2 = 100$	$d_2 = 80$	8000	0,01000
120	160	$a_3 = 140$	$d_3 = 60$	8400	0,00714
160	200	$a_4 = 180$	$d_4 = 40$	7200	0,00555

Если учесть время, проведенное на каждом уровне цены, становится видно, что средняя цена должна быть ниже \$ 1,20. Правильную среднюю цену можно рассчитать при условии, что известно количество дней в каждом интервале, используя средневзвешенное значение цены

$$W = \frac{a_1 d_1 + a_2 d_2 + a_3 d_3 + a_4 d_4}{d_1 + d_2 + d_3 + d_4}$$

и соответствующий интервал

$$W = \frac{6000 + 8000 + 8400 + 7200}{280}$$

$$W = 105,71.$$

Этот результат может изменяться в зависимости от количества используемых временных интервалов, однако он дает лучшее представление о правильной средней цене. Существуют две других средних величины, для которых время является важным элементом, — это среднее геометрическое и среднее гармоническое.

Среднее геометрическое

Среднее геометрическое представляет функцию роста, в которой изменение цены от 50 до 100 столь же важно, как и изменение от 100 до 200. Если существует n цен, $a_1, a_2, a_3 \dots a_n$, то геометрическим средним является n корень из произведения цен

$$G = (a_1 \times a_2 \times a_3 \times \dots \times a_n)^{\frac{1}{n}}$$

или

$$\text{product}(a_1, a_2, a_3, \dots, a_n)^{\frac{1}{n}}.$$

Чтобы решить это математически, а не используя электронную таблицу, нужно преобразовать приведенное выше уравнение в любую из двух форм:

$$\ln(G) = \frac{\ln(a_1) + \ln(a_2) + \dots + \ln(a_n)}{n}$$

или

$$\ln(G) = \frac{\ln(a_1 \times a_2 \times a_3 \times \dots \times a_n)}{n}.$$

Эти два решения эквивалентны. Член $\ln$ является натуральным логарифмом. (Заметьте, что в некоторых программах функция $\log$ фактически представляет собой $\ln$.) Используя уровни цен из табл. 2.1, записываем

$$\ln(G) = \frac{\ln(40) + \ln(80) + \ln(120) + \ln(160) + \ln(200)}{5}.$$

Игнорируя временные интервалы, подставляем в первое уравнение:

$$\ln(G) = \frac{3,689 + 4,382 + 4,787 + 5,075 + 5,298}{5}.$$

Следовательно:

$$\ln(G) = 4,6462$$

$$G = 104,19$$

Если среднее арифметическое, взвешенное по времени, составляет 105,71, то среднее геометрическое дает 104,19.

Среднее геометрическое имеет преимущества применительно к экономике и ценам. Классический пример — сравнение десятикратного повышения цены со 100 до 1000 с падением до одной десятой со 100 до 10. Среднее арифметическое двух величин (10 и 1000) равно 505, а среднее геометрическое дает

$$G = (10 \times 1000)^{\frac{1}{2}} = 100$$

и показывает относительное распределение цен как функцию сопоставимого роста. Благодаря этому свойству геометрическое среднее лучше подходит для усреднения коэффициентов, которые могут представляться в виде как дроби, так и процента.

Среднее квадратичное

Среднее квадратичное чаще всего используется для оценки погрешности. Оно рассчитывается следующим образом:

$$Q = \sqrt{\frac{\sum a^2}{N}}.$$

Среднее квадратичное представляет собой квадратный корень из среднего арифметического квадратов величин. Лучше всего оно известно как основа для расчета стандартного отклонения. Мы поговорим об этом далее в этой главе в разделе «Моменты распределения: дисперсия, асимметрия и эксцесс».

Среднее гармоническое

Среднее гармоническое — еще одно взвешенное по времени среднее, но оно не тяготеет к более высоким или более низким величинам, как среднее геометрическое. В качестве простого примера рассмотрим среднюю скорость авто-

мобиля, проезжающего 4 мили со скоростью 20 миль в час, а затем 4 мили со скоростью 30 миль в час. Среднее арифметическое дало бы 25 миль в час, без учета того, что 12 минут машина ехала со скоростью 20 миль в час, а 8 минут — 30 миль в час. Средневзвешенное значение составило бы

$$W = \frac{(12 \times 20) + (8 \times 30)}{12 + 8} = 24.$$

Среднее гармоническое равно

$$\frac{1}{H} = \frac{\frac{1}{a_1} + \frac{1}{a_2} + \dots + \frac{1}{a_n}}{n},$$

что можно также выразить как

$$H = \frac{n}{\sum_{i=1}^n \left(\frac{1}{a_i} \right)}.$$

Для двух-трех значений можно использовать более простую форму:

$$H_2 = \frac{2ab}{a+b}$$

$$H_3 = \frac{3abc}{ab+ac+bc}$$

Это позволяет видеть некоторую закономерность в решении. При скоростях 20 и 30 миль в час решение представляется как

$$H_2 = \frac{2 \times 20 \times 30}{20 + 30} = 24,$$

что аналогично средневзвешенному значению. Рассмотрим еще раз первоначальное множество чисел, применив базовую форму среднего гармонического:

$$\frac{1}{H} = \frac{\frac{1}{40} + \frac{1}{80} + \frac{1}{120} + \frac{1}{160} + \frac{1}{200}}{5},$$

$$H = 87,59 = \frac{0,5708}{5} = 0,01142.$$

Мы могли бы применить среднее гармоническое к колебаниям цен, где первое колебание равно 20 пунктам за 12 дней, а второе колебание составляет 30 пунктов за 8 дней.

Распределение цен

Измерение распределения очень важно, поскольку дает общее представление о том, чего ожидать. Мы не можем знать, каким будет завтра торговый диапазон S&P, но если текущая цена равна 1200, то мы можем утверждать с высокой уверенностью, что в этом году она будет находиться между 900 и 1500, и с меньшей уверенностью, что она будет между 1100 и 1300. Еще меньше уверенности у нас в том, что она ограничится коридором между 1150 и 1250, и нет практически никакого шанса угадать диапазон точно. Нижеследующие методы измерения распределения позволят вам определять вероятность (т. е. доверительный уровень) наступления события.

Во всех статистических примерах, включенных в эту книгу, мы будем использовать в качестве выборочных данных ограниченное количество цен или — в некоторых случаях — отдельные торговые прибыли и убытки. Мы будем измерять характеристики выборки, находя форму распределения, решая, как результаты малой выборки соотносятся с большой или насколько две выборки подобны друг другу. Все эти измерения показывают, что малые выборки менее надежны, но их можно использовать, если вы понимаете величину ошибки или различие в форме распределения по сравнению с ожидаемым распределением большой выборки.

Плотность распределения

Плотность распределения (называемая также *гистограммой*) концептуально весьма проста, но при этом может давать хорошую картину характеристик данных. Теоретически мы ожидаем, что цены на биржевые товары будут больше времени находиться на низких уровнях, повышаясь лишь ненадолго. Эта модель показана на рис. 2.2 (цены на пшеницу за последние 25 лет). Чаше всего цена



Рис. 2.2. Цены на пшеницу, 1985–2010 гг.

находится там, где спрос и предложение сбалансированы, это называется *равновесием*. Когда предложение недостаточно или возникает неожиданный спрос, цена на короткое время повышается, пока или спрос не удовлетворится (что может произойти, если цены слишком высоки), или предложение не увеличится до уровня спроса. Обычно слева есть небольшой хвост, где в течение периодов высокого предложения цены иногда опускаются ниже себестоимости.

Чтобы рассчитать плотность распределения для данных, разделенных на 20 сегментов, найдем самую высокую и самую низкую из рассматриваемых цен и разделим разность на 19, чтобы получить размер одного сегмента. Затем, начиная с самой низкой цены, прибавляем размер сегмента, чтобы получить второе значение, прибавляем размер сегмента ко второму значению и получаем величину третьего и так далее. В конце получится 20 сегментов, которые начинаются на самой низкой цене и заканчиваются на самой высокой. После этого можно подсчитать количество цен, попадающих в каждый сегмент, что является задачей почти невыполнимой, или использовать для этого электронную таблицу. В Excel нужно пройти в Данные/Анализ данных/Гистограмма и ввести диапазон сегментов (который нужно задать заранее) и данные для анализа, затем выбрать незаполненное место в таблице для вывода результатов (место справа от сегментов подойдет) и нажать ОК. Плотность распределения будет показана немедленно. Затем можно представить результаты графически, как показано на рис. 2.3.

Плотность распределения показывает, что чаще всего цена располагалась между \$3,50 и \$4,00 за бушель, но весь активный диапазон простирался от \$2,50 до \$5,00. Справа хвост простирается до почти \$10 за бушель. Это означает, что в распределении цены присутствует *толстый хвост*. При нормальном распределении не было бы цен выше \$6. Отсутствие ценовых данных ниже \$2,50 объясняется порогом себестоимости. Ниже этой цены фермеры отказываются продавать себе в убыток, однако у американского правительства есть программа поддержки цен, которая гарантирует фермерам минимальную прибыль.

Плотность распределения цен на пшеницу можно также рассматривать с учетом инфляции или изменения курса доллара США. Мы вернемся к этому в конце данной главы.



Рис. 2.3. Плотность распределения цен на пшеницу, демонстрирующая хвост справа

Краткосрочное распределение

Аналогичные плотности распределения встречаются и при изучении более коротких временных интервалов, хотя чем временной интервал короче, тем картина хаотичнее. Взяв цены на пшеницу за 2007 календарный год (рис. 2.4), мы видим устойчивое восходящее движение в середине года, за которым следует широкодиапазонное боковое движение на более высоком уровне. Однако плотность распределения на рис. 2.5 показывает модель, подобную долгосрочному распределению, где наиболее распространенные значения находятся на низком уровне и присутствует толстый хвост справа. Если бы мы выбрали несколько месяцев непосредственно перед тем, как цены достигли максимума в сентябре 2007 г., на графике ценовой пик сдвинулся бы вправо и появился бы толстый хвост слева. Для биржевых товаров это означает период ценовой нестабильности и ожидания падения цен.

Следует ожидать, что распределение цен физических товаров, таких как сельскохозяйственная продукция, металлы и энергоносители, будет смещено



Рис. 2.4. Пшеница, дневные цены, 2007 г.

Распределение цен на пшеницу, 2007 г.

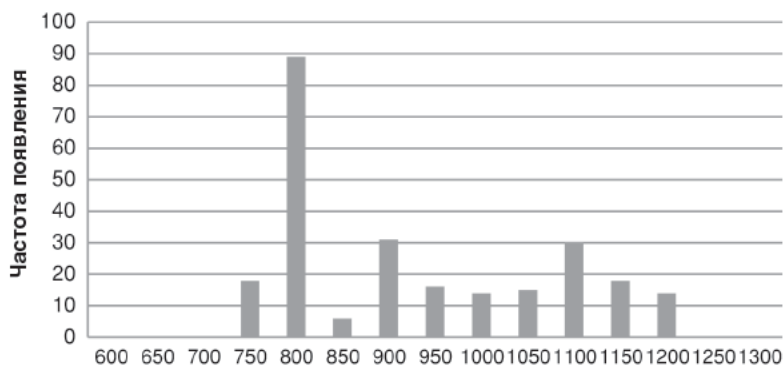


Рис. 2.5. Плотность распределения цен на пшеницу, интервалы \$0,50, в течение 2007 г.

влево (низкие цены встречаются чаще) и иметь *длинный хвост* в области высоких цен в правой части графика. Это происходит потому, что цены держатся на относительно более высоких уровнях в течение лишь короткого времени, пока существует дисбаланс между спросом и предложением. История фондового рынка показывает, что у акций не может до бесконечности поддерживаться исключительно высокое отношение «цена/прибыль» (P/E). Однако период корректировки здесь может растягиваться на многие годы в отличие от сельскохозяйственных продуктов, циклы которых начинаются заново каждый год. При рассмотрении более коротких ценовых периодов модели, не соответствующее стандартному распределению, можно считать переходными. Читатели, желающие глубже исследовать эту тему, должны прочитать главу 18, особенно разделы «Распределение цен» и «Профиль рынка Стидлмайера».

Меры *среднего значения*, описанные в предыдущем разделе, используются для описания траектории и экстремумов движения цены в плотности распределения. Общее соотношение между тремя главными мерами в условиях, когда распределение не идеально симметрично, выглядит следующим образом:

$$\text{Среднее арифметическое} > \text{среднее геометрическое} > \\ > \text{среднее гармоническое.}$$

Медиана и мода

Для определения параметров распределения часто используются еще два измерения — медиана и мода. *Медиана*, или «середина», полезна для нахождения «центра» данных — если данные упорядочены, это то значение, которое находится в середине. Медиана позволяет исключить влияние экстремумов, которые могут исказить среднее арифметическое. Недостатком ее является то, что для отыскания средней точки нужно упорядочить все данные. Медиана предпочтительнее среднего, кроме тех случаев, когда используется очень небольшое количество данных.

Мода — это наиболее часто встречающееся значение. На рисунке 2.5 модой является самый высокий столбик плотности распределения в сегменте 800.

При нормальном распределении ценового ряда мода, среднее и медиана равны, однако чем больше нарушается симметрия данных, тем дальше расходятся эти параметры. Их взаимоотношение выглядит следующим образом:

$$\text{Среднее} > \text{медиана} > \text{мода.}$$

Нормальное распределение обычно называют *колоколообразной кривой*, где значения располагаются поровну с обеих сторон от среднего. В большинстве случаев при работе с ценовыми данными и результатами торговли распределение смещается вправо (к более высоким ценам или более высокой торговой прибыли) и сглаживается или обрезается слева (более низкие цены или торговые убытки). Если бы вам потребовалось построить график распределения торговых прибылей и убытков для системы следования за трендом с фиксированным стоп-лоссом, вы получили бы прибыли, варьирующие от нуля до очень больших величин, а вот убытки были бы теоретически ограничены размером

стоп-лосса. Асимметричные распределения будут важны, когда позднее в этой главе мы займемся измерением вероятности. В торговой среде «нормальных» распределений не бывает.

Характеристики основных методов усреднения

Каждый метод усреднения обладает своим уникальным смыслом и полезностью. Ниже вкратце обобщаются их основные характеристики.

На *среднее арифметическое* одинаково влияет каждый элемент данных, но оно чувствительно к экстремумам больше, чем другие методы. Его легко рассчитывать и можно использовать в алгебраических формулах.

Среднее геометрическое придает экстремальным отклонениям меньше веса, чем среднее арифметическое, и оно наиболее важно при использовании данных, представляющих относительные величины (коэффициенты) или темпы изменения. Его нельзя использовать с отрицательными числами, но можно включать в алгебраические формулы.

Среднее гармоническое лучше всего применимо к временным изменениям и, наряду с средним геометрическим, используется в экономике для анализа цен. Оно рассчитывается труднее и потому менее популярно, чем любое другое среднее, но его также можно включать в алгебраические формулы.

Мода — это самое распространенное значение, определяемое только плотностью распределения. Это точка наибольшей концентрации, она указывает типичное значение в разумно большой выборке. Чтобы найти моду в неупорядоченном множестве данных, таком как цены, требуется очень много времени. Использовать ее в алгебраических формулах нельзя.

Медиана представляет собой срединное значение. Она наиболее полезна, когда необходимо найти центр неполного множества. Она нечувствительна к экстремальным колебаниям, и ее просто найти, но для этого требуется упорядочение данных, что замедляет вычисление. Хотя у нее есть некоторые арифметические свойства, в чистом виде использовать ее в расчетах нельзя.

Моменты распределения: дисперсия, асимметрия и эксцесс

Моменты распределения описывают форму распределения точек данных, т. е. их расположение вокруг среднего. Существует четыре таких момента: *среднее*, *дисперсия*, *асимметрия* и *эксцесс*, причем каждый описывает отдельный аспект формы распределения. Попросту говоря, среднее представляет собой центр, или среднее значение, дисперсия — это удаленность отдельных точек от среднего, асимметрия — это смещение распределения влево или вправо относительно среднего, а эксцесс — острровершинность распределения. Мы уже описали среднее, поэтому начнем со второго момента.

В последующих расчетах мы будем использовать обозначение $\bar{P}$ для указания среднего значения ряда, состоящего из n цен. Прописная буква P обозначает совокупность всех цен, а строчная p — отдельные цены.

$$\bar{P} = \frac{\sum_{i=1}^n P_i}{n}.$$

Основным измерителем распределения является *среднее отклонение* (mean deviation — MD), его рассчитывают относительно центра распределения, например относительно среднего арифметического.

$$MD = \frac{\sum_{i=1}^n |P_i - \bar{P}|}{n}.$$

Таким образом, MD равно средней величине разностей между каждой ценой и средним арифметическим этих цен или какой-то другой характеристикой центра распределения, при этом все разности должны рассматриваться как положительные числа. В книге эта формула встречается очень часто.

Дисперсия (второй момент)

Дисперсия (Var), которая очень похожа на среднее отклонение — лучшую оценку дисперсии, — используется как основа для многих других вычислений. Она равна

$$Var = \frac{\sum_{i=1}^n (P_i - \bar{P})^2}{n - 1}.$$

Обратите внимание на то, что дисперсия равна квадрату стандартного отклонения, $var = s^2 = \sigma^2$, это один из наиболее распространенных статистических показателей. В Excel дисперсия записывается как функция *var (list)*, а в TradeStation EasyLanguage это *variance (series, n)*.

Стандартное отклонение (s), обозначаемое чаще всего буквой σ (сигма), является особой формой измерения среднего отклонения от среднего, в которой используется среднеквадратичное отклонение

$$\sigma = \sqrt{\frac{\sum_{i=1}^n (P_i - \bar{P})^2}{n}},$$

где разности между отдельными ценами и средней ценой возводятся в квадрат, чтобы повысить значимость экстремумов, а затем результат приводится в норму путем извлечения квадратного корня. Этот популярный показатель, часто используемый в настоящей книге, определяется в Excel с помощью функции *Stdevp*, а в TradeStation — функции *StdDev (price, n)*, где n — количество цен.

Стандартное отклонение является самым популярным измерителем дисперсии данных. Одно стандартное отклонение от среднего представляет совокупность приблизительно 68% данных, два стандартных отклонения от среднего включают 95,5% всех данных, а три стандартных отклонения охватывают 99,7%, т. е. почти все данные. И хотя гарантировать включение всех данных невозможно, в случае нормального распределения вы можете использовать 3,5 стандартных отклонения, чтобы включить 100% данных. Эти величины

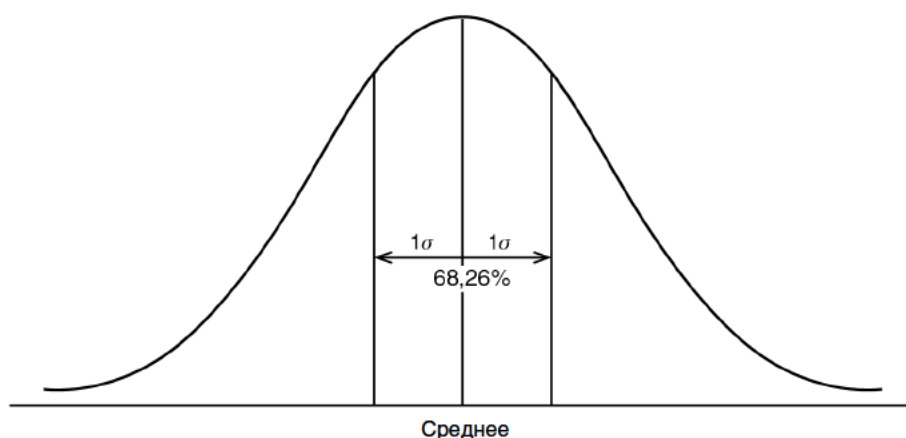


Рис. 2.6. Нормальное распределение и площадь в процентах, охватываемая одним стандартным отклонением от среднего арифметического

представляют группировки идеально *нормального* множества данных. Они показаны на рис. 2.6.

Асимметрия (третий момент)

В большинстве случаев, однако, ценовые данные имеют ненормальное распределение. У физических биржевых товаров, таких как золото, зерновые, энергоносители и даже процентные ставки (выраженные как доходность), цены находятся больше времени на низких уровнях и намного меньше времени на максимумах. Например, достигнув максимума \$ 800 за унцию в январе 1980 г., золото продержалось на этом уровне всего один день и оставалось в диапазоне \$ 250–400 за унцию в течение большей части следующих 20 лет. Если взять в качестве среднего \$ 325, то распределение цены просто не может быть симметричным. Если одно стандартное отклонение равно \$ 140, то нормальное распределение показало бы высокую вероятность падения цены до \$ 185, а это крайне маловероятный сценарий. Асимметрия проявляется очевиднее всего на сельскохозяйственных рынках, где в один год дефицит сои или кофе может очень сильно взвинтить цены, а на следующий год нормальный урожай возвратит цены на прежние уровни.

Взаимосвязь цены и времени, при которой рынки дольше находятся на более низких уровнях, можно определить как *асимметрию* — величину отклонения от симметричного распределения, заставляющего кривую выглядеть короче слева и длиннее справа (где более высокие цены). Удлиненную сторону называют *хвостом*, и если более длинный хвост находится справа, то говорят о *положительной асимметрии*. В случае *отрицательной асимметрии* хвост находится слева. Это видно на рис. 2.7.

При абсолютно нормальном распределении среднее, медиана и мода совпадают. Когда цены демонстрируют положительную асимметрию, что типично для периода более высоких цен, среднее подвергается наибольшему измене-

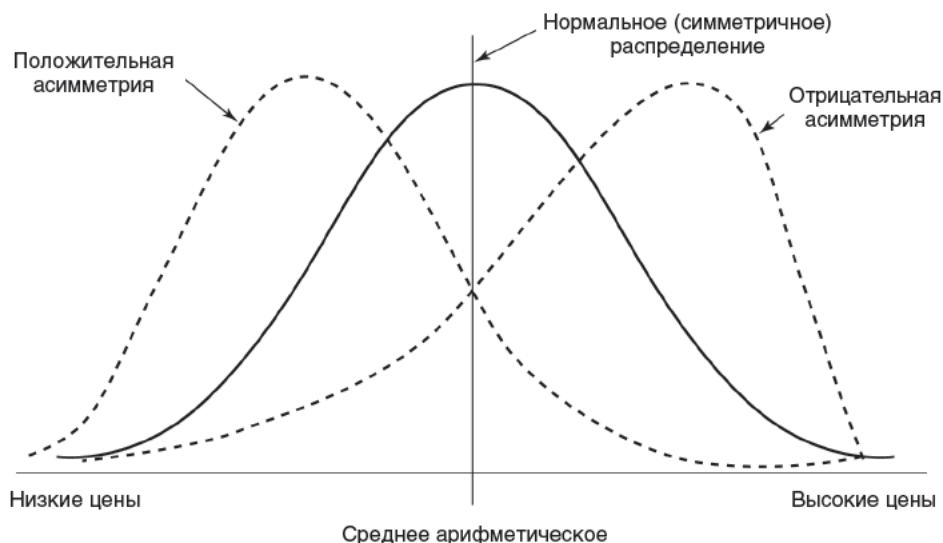


Рис. 2.7. Асимметрия. Почти все распределения цен имеют положительную асимметрию с более длинным хвостом справа, где цены выше

нию, мода — наименьшему, а медиана оказывается где-то посередине. Хорошее представление об асимметрии дает разность между средним и модой с учетом дисперсии в виде стандартного отклонения распределения.:

$$(\text{Асимметрия}) S_K = \frac{\text{Среднее} - \text{мода}}{\text{Стандартное отклонение}}.$$

В умеренно асимметричном распределении расстояние между средним и модой оказывается в три раза больше разности между средним и медианой. Эту взаимосвязь можно также записать следующим образом:

$$(\text{Асимметрия}) S_K = \frac{3 \times (\text{среднее} - \text{медиана})}{\text{Стандартное отклонение}}.$$

Чтобы продемонстрировать подобие второго и третьего моментов (дисперсия и асимметрия), чаще используется следующая формула:

$$S_K = \frac{\sum_{i=1}^n (p_i - \bar{P})^3}{(n-1)\sigma^3},$$

где n — количество цен в распределении, а σ — стандартное отклонение цен.

Функции для расчета асимметрии есть как в Excel, так и в TradeStation.

Преобразования

Иногда асимметрию ряда данных можно устранить с помощью *преобразования*. Симметрия ценовых данных может нарушаться конкретными факторами.

Например, если есть три случая удвоения цены и $1/9$ случая утроения, то первоначальные данные можно преобразовать в нормальное распределение путем извлечения квадратного корня из каждого элемента данных. Характеристики ценовых данных часто имеют логарифмический, степенной или квадратно-корневой характер.

Чтобы рассчитать уровень вероятности распределения, основываясь на асимметричном распределении цены, можно преобразовать нормальную вероятность в эквивалентную экспоненциальную вероятность (P_E), используя

$$P_E = \frac{\bar{X} \log_{10} \left(\frac{1}{1-P} \right)}{\log_{10} e},$$

где $\bar{X}$ — среднее всех цен;
 P — нормальная вероятность;
 $\log_{10} e = 0,434294482$.

В то время как нормальная вероятность P преуменьшает вероятность возникновения события в распределении цены, экспоненциальное распределение P_E преувеличивает ее. Когда возможно, лучше использовать точные вычисления, однако при расчете риска может быть целесообразнее немного преувеличивать ожидаемый риск.

Асимметрия в распределениях на различных относительных уровнях цен

Поскольку минимальные уровни цен большинства биржевых товаров определяются себестоимостью производства, распределения цен демонстрируют явное сопротивление падению ниже этих порогов. Это способствует возникновению положительной асимметрии на рынках соответствующих товаров. Учитывая, что цены находятся на необычно высоких уровнях лишь в течение коротких периодов, они могут быть волатильными, а вызываемая ими отрицательная асимметрия может характеризоваться как неустойчивая. Между очень высокими и очень низкими уровнями цен мы можем найти плотность распределения, которая выглядит нормальной. На рисунке 2.8 показано изменение в распределении цен в течение, скажем, 20 дней, когда цены резко идут вверх. Среднее показывает центры распределения по мере того, как оно изменяется, сдвигаясь от положительной к отрицательной асимметрии. Этот пример показывает, что нормальное распределение не подходит для всех случаев анализа цен, и что логарифмическое, экспоненциальное или степенное распределение подходит лучше всего для долгосрочного анализа.

Эксцесс (четвертый момент)

Наконец, последнее измерение, *эксцесс*, необходимо для описания формы ценового распределения. На рисунке 2.9 показаны *эксцессы*, отражающие заостренность или сглаженность распределения. Этот измеритель хорош для беспристрастной оценки того, формируют ли цены тренд или движутся вбок. Если вы

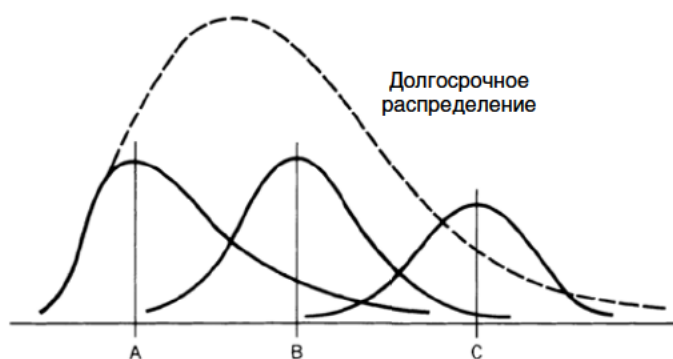


Рис. 2.8. Изменение распределения на различных уровнях цен. А, В и С — последовательно увеличивающиеся средние значения трех краткосрочных распределений. Они показывают, как распределение изменяется, сдвигаясь от положительной к отрицательной асимметрии

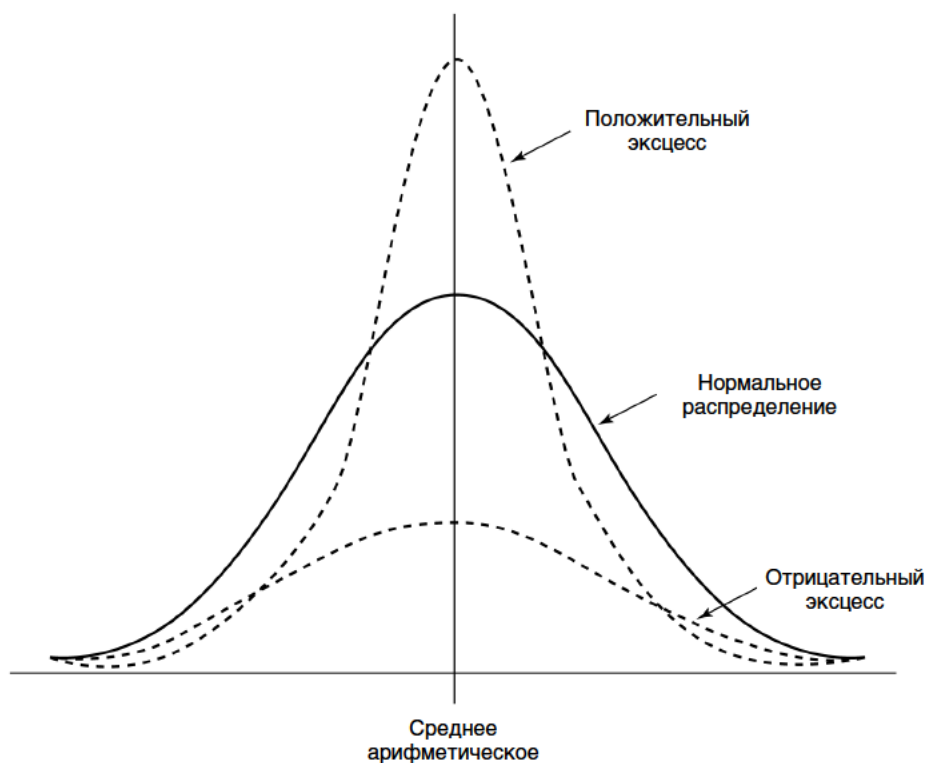


Рис. 2.9. Эксцесс. Положительный эксцесс образуется, когда пик распределения выше нормального, что типично для бокового рынка. Отрицательный эксцесс, выглядящий как более плоское распределение, формируется, когда рынок находится в тренде

видите, что цены равномерно повышаются, распределение будет более плоским и охватит более широкий диапазон. Это называется *отрицательным эксцессом*. Если цены идут вбок, образуется скопление вокруг среднего, и мы имеем *положительный эксцесс*. В профиле рынка Стидлмайера, описываемом в главе 18, используется концепция эксцесса с плотностью распределения, накапливаемой динамически за счет учета изменения цены в реальном времени.

Следуя той же форме, что и в случае третьего момента, асимметрии, эксцесс можно рассчитать как

$$K = \frac{\sum_{i=1}^n (p_i - \bar{P})^4}{(n-1)\sigma^4}.$$

По-другому эксцесс можно рассчитать так:

$$K = \frac{n(n+1)}{(n-1)(n-2)(n-3)} \sum \left(\frac{p_i - \bar{P}}{\sigma} \right)^4 - \frac{3(n-1)^2}{(n-2)(n-3)},$$

где n — количество цен в распределении;

p_i — отдельные цены;

$\bar{P}$ — среднее n цен;

σ — стандартное отклонение цен.

Чаше всего используется *модифицированный показатель эксцесса* (excess kurtosis, обозначается KE), позволяющий видеть ненормальные распределения лучше. $KE = K - 3$, потому что нормальная величина эксцесса равна 3.

Эксцесс также полезен при анализе тестов систем. Когда вы рассчитываете эксцессы дневной доходности, они должны быть несколько выше нормального, если система прибыльна. Однако, если эксцесс оказывается выше 7–8, может оказаться, что метод торговли чрезмерно подогнан. Высокий эксцесс означает, что имеется слишком много прибыльных сделок одинакового размера, что вряд ли возможно в реальной торговле. Высокое значение эксцесса должно сразу же вызывать подозрения.

Выбор между плотностью распределения и стандартным отклонением

Плотность распределения важна, поскольку стандартное отклонение не работает в случае асимметричных распределений, наиболее характерных для большинства ценовых данных. Вернемся для примера к гистограмме пшеницы. Средняя цена за предшествующие 25 лет составляла \$3,62 при стандартном отклонении цен \$1,16. В этом случае одно стандартное отклонение влево от среднего дает \$2,46, но в этом сегменте нет никаких данных. На правой же стороне, на удалении 3,5 стандартного отклонения, которое должно охватывать 100% данных, находится цена \$7,68, но она значительно ниже фактической максимальной цены.

Значит, использование стандартного отклонения может давать неверный результат на обоих концах распределения очень асимметричных данных,

в то время как плотность распределения дает очень ясную и полезную картину. Если бы мы захотели, основываясь на плотности распределения, узнать цену на 10%-ном и 90%-ном уровнях вероятности, то нам нужно было бы отсортировать все данные от минимума до максимума. И если было 300 месячных точек данных, то 10%-ный уровень находился бы в точке 30, а 90%-ный уровень — в точке 271. Медианная цена располагалась бы в точке 151. Это показано на рис. 2.10.

В случае длинного хвоста справа и плотность распределения, и стандартное отклонение подразумевают, что следует ожидать больших движений. Когда распределение очень симметрично, беспокоиться особо не о чем. В отношении рынков, уже переживших экстремальные движения, ни один из этих методов не скажет, какого размера экстремальное движение может произойти. Несомненно, если пройдет достаточное время, мы увидим прибыли и убытки, которые будут больше тех, что наблюдались в прошлом, возможно, намного больше.

Автокорреляция

Сериальная корреляция, или *автокорреляция*, означает, что данным присуще некоторое постоянство, т. е. будущие данные можно предсказать (до некоторой степени) на основе прошлых данных. Существование этого качества может говорить о наличии трендов. Автокорреляцию можно легко найти, если поместить данные в столбец А электронной таблицы, а затем скопировать их в столбец В, сместив вниз на одну строку. Затем определяется корреляция столбцов А и В. Дополнительные корреляции можно рассчитать, смещая столбец В вниз на 2, 3 или 4 строки. Так можно обнаружить существование цикла.

Формальным способом нахождения автокорреляции является использование критерия Дарбина — Уотсона, позволяющего получить *d*-статистику. Этот под-

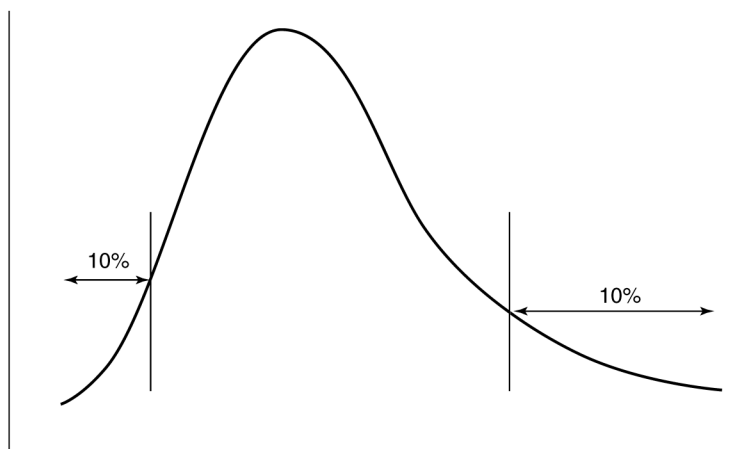


Рис. 2.10. Отсечение 10% от каждого конца плотности распределения.

Из-за высокой плотности низких цен эта зона кажется узкой, в то время как менее плотные данные в области высоких цен зрительно занимают большее пространство

ход измеряет изменение остатков (e), т. е. разность между N точек данных и их средним значением.

$$e_t = r_t - \frac{\sum_{t-N+1}^t r_i}{N}.$$

$$d = \frac{\sum_{t-N+1}^t (e_i - e_{i-1})^2}{\sum_{t-N+1}^t e_i^2}.$$

Величина d всегда находится между 0 и 4. Если $d = 2$, никакой автокорреляции нет. Если d существенно меньше 2, существует положительная автокорреляция, однако при d ниже 1 сходство остатков превышает разумный уровень. Чем больше d превышает 2, тем более отрицательной представляется автокорреляция.

Наличие *положительной автокорреляции*, или сериальной корреляции, означает, что существует хороший шанс на повторение рассматриваемых событий в будущем.

Вероятность получения прибыли

Сомнение заставляет испытывать
неудобство, но уверенность делает человека
смешным.

Китайская пословица

Если бы годовая доходность фондового рынка за последние 50 лет имела бы нормальное распределение (рис. 2.6), то среднее составило бы приблизительно 8%, а одно стандартное отклонение — 16%. В любом отдельно взятом году мы можем ожидать доходность на уровне 8%, однако существует 32%-ная вероятность, что она будет или больше 24% (среднее плюс одно стандартное отклонение), или меньше –8% (среднее минус одно стандартное отклонение). Если вы захотели бы узнать вероятность получения прибыли в 20% или больше, вам сначала пришлось бы пересчитать формулу

$$\text{Вероятность достижения цели} = \frac{\text{Цель} - \text{среднее}}{\text{Стандартное отклонение}}.$$

Если цель равна 20%, то

$$\text{Вероятность} = \frac{20\% - 8\%}{16\%} - 0,75.$$

Таблица A1.1 в приложении 1 содержит вероятности для нормальных кривых. Поиск стандартного отклонения 0,75 дает 27,34%, это совокупность 54,68% данных. После этого выше цели 20% оказывается половина оставшихся данных, или 22,66%.

Автоматический расчет вероятности

Неудобно искать значения вероятности в таблице, когда вы работаете с электронной таблицей или компьютерной программой, но вероятности понять легче, чем значения стандартного отклонения. Вы можете рассчитать область под кривой, которая соответствует конкретному значению z (стандартное отклонение), используя следующую аппроксимацию¹.

Допустим, что $z' = |z|$, абсолютному значению z . Тогда

$$r = 1 + z' \times (c_1 + z' \times (c_2 + z' \times (c_3 + z' \times (c_4 + z' \times c_5)))) ,$$

где $c_1 = 0,049867347$;
 $c_2 = 0,0211410061$;
 $c_3 = 0,032776263$;
 $c_4 = 0,0000380036$;
 $c_5 = 0,0000488906$;
 $c_6 = 0,000005383$.

Отсюда вероятность P того, что доходность будет равна или выше ожидаемой доходности, составляет

$$P = 0,5 \times e^{[\ln(r) \times (-16)]} .$$

Используя пример, в котором стандартное отклонение $z = 0,75$, производим расчет

$$\begin{aligned} r = 1 + 0,75 \times (0,049867347 + 0,75 \times [0,0211410061 + \\ 0,75 \times [0,0000380036 + 0,75 \times (0,0327762632 + \\ + 0,75 \times (0,0000488906 + 0,75 \times [0,000005383])]])]) \\ r = 1,0507. \end{aligned}$$

Подставив значение r в уравнение P , получаем

$$P = 0,5 \times e^{[\ln(1,0507) \times (-16)]} = 0,226627 .$$

Следовательно, существует 22,7%-ная вероятность того, что значение превысит 0,75 стандартного отклонения (т. е. окажется на одной из сторон распределения за пределами 0,75). Вероятность того, что значение окажется внутри зоны, ограниченной $\pm 0,75$ стандартного отклонения, равно $1 - (2 \times 0,2266) = 0,5468$, или 54,68%. Это же значение мы находим в табл. А. 1.1 приложения 1.

Те, кто пользуется Excel, могут найти ответ с помощью функции *normdist* (p , *mean*, *stdev*, *cumulative*), где

p — текущая цена или значение;

mean — среднее ряда p ;

stdev — стандартное отклонение ряда p , и *cumulative* имеет значение true, если вы хотите получить z .

¹ Stephen J. Brown and Mark P. Kritzman, *Quantitative Methods for Financial Analysis*, 2nd ed. (Dow Jones-Irwin, 1990), 238–241.

Отсюда результат *normdist* (35,20,5, true) равен 0,99 865, или вероятность 99,8%, а если *cumulative* имеет значение false, то результат будет равен 0,000866.

Стандартная ошибка

В процессе разработки и тестирования торговой системы мы хотим знать, соответствуют ли результаты, которые мы видим, ожидаемым. Ответ всегда зависит от размера выборки данных и величины дисперсии, типичной для данных в течение соответствующего периода.

Существует описательный метод измерения ошибки, получивший название *среднеквадратичная ошибка*, или *стандартная ошибка* (standard error — SE). В нем используется дисперсия, позволяющая рассчитать ошибку по распределению данных при наличии нескольких выборок данных. В этом тесте определяется, как средние значения каждой выборки отличаются от среднего значения всей совокупности данных, т. е. речь идет об *однородности* данных.

$$SE = \sqrt{\frac{Var}{n}},$$

где *Var* = дисперсия средних значений выборок;

n = количество точек данных в средних значениях выборок.

Когда мы говорим *средние значения выборок*, то имеем в виду, что выборки данных осуществляются неоднократно, каждая включает *n* точек данных и для нахождения дисперсии используются средние значения этих выборок. Однако в большинстве случаев мы будем использовать единственный ряд данных, рассчитывая дисперсию, как было показано в этой главе ранее.

t-статистика и степени свободы

Когда в распределении используется меньше цен или сделок, можно ожидать, что форма кривой будет более изменчивой. Например, пик распределения может оказаться ниже, а хвосты выше. Чтобы измерить, насколько распределение выборки из меньшего множества близко к нормальному распределению (большой выборки данных), можно использовать *t-статистику* (ее также называют *t-критерием Стьюдента*, который разработал Уильям Госсет). *t*-критерий рассчитывается в зависимости от степени свободы (*df*), которая равна *n* – 1, где *n* представляет собой размер выборки, т. е. количество цен, используемых в распределении.

$$t = \frac{\text{Среднее изменение цены}}{\text{Стандартное отклонение}} \times \sqrt{n}.$$

Чем больше данных в выборке, тем надежнее результаты. Мы можем получить общее представление о форме распределения, взглянув на несколько значений *t* в табл. 2.2, где указаны значения *t*, соответствующие верхним областям хвоста 0,10, 0,05, 0,025, 0,01 и 0,005. Таблица показывает, что по мере увели-

чения размера выборки n значение t приближается к величинам, характерным для нормальных значений стандартного отклонения в областях хвоста.

Наиболее важные значения t можно найти в табл. A1.2 « t -распределение» приложения 1. Столбец 0,10 дает 90%-ный доверительный уровень, 0,05 — 95%-ый, а 0,005 — 99,5%-ный. Например, если в выборке 20 цен и мы хотим иметь вероятность верхнего хвоста 0,025, то величина t должна быть равна 2,086. Для меньших выборок величина t должна быть больше, чтобы обеспечить тот же доверительный уровень.

При тестировании торговой системы степень свободы может определяться количеством сделок, предлагаемых стратегией. Когда сделок немного, результаты могут не отражать того, что вы получите при более продолжительной торговле. Тестируя стратегию, вы находите похожую взаимосвязь между числом сделок и количеством параметров, или переменных, используемых в стратегии. Чем больше переменных, тем больше требуется сделок, чтобы создать приемлемый доверительный уровень.

t -критерий для двух выборок

Таблица 2.2. Значения t , соответствующие вероятности появления верхнего хвоста 0,025

Степень свободы (df)	Величина t
1	12,706
10	2,228
20	2,086
30	2,042
120	1,980
Нормальная	1,960

У вас может возникнуть необходимость сравнения двух периодов данных, чтобы решить, не произошло ли значительных изменений в поведении цены. Некоторые аналитики используют этот прием для устранения противоречивых данных, однако характеристики цен и экономические данные изменяются в результате эволюционного процесса, и системная торговля должна приспосабливаться к этим изменениям. Этот тест лучше всего применять к результатам торговли, когда нужно понять, насколько устойчиво работает стратегия. Это делается с помощью t -критерия для двух выборок:

$$t = \frac{\bar{P}_1 - \bar{P}_2}{\sqrt{\frac{\text{var}_1^2}{n_1} + \frac{\text{var}_2^2}{n_2}}},$$

где $\bar{P}_1$ и $\bar{P}_2$ — средние значения цен периодов 1 и 2;
 var_1 и var_2 — дисперсия цен периодов 1 и 2;
 n_1 и n_2 — количество цен в периодах 1 и 2,

а два сравниваемых периода являются взаимно исключаящими. Степени свободы df , необходимые для определения доверительных уровней в табл. А1.2, можно рассчитать, используя аппроксимацию Саттертвейта, где s — стандартное отклонение значений данных:

$$df = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} \right)}{\frac{\left(\frac{s_1^2}{n_1} \right)}{(n_1 - 1)} + \frac{\left(\frac{s_2^2}{n_2} \right)}{(n_2 - 1)}} .$$

Используя t -критерий для определения стабильности прибылей и убытков, генерируемых торговой системой, замените элементы данных на чистую доходность каждой сделки, количество элементов данных количеством сделок и рассчитайте все остальные значения, используя доходность, а не цены.

Нормализация риска и доходности

Чтобы сравнивать один метод торговли с другим, необходимо нормализовать как тесты, так и параметры, используемые для оценки. Если одна система имеет совокупную доходность 50%, а другая 250%, мы не можем определить, какая из них лучше, если не знаем продолжительности тестов и волатильности доходности или риска. Если 50%-ная доходность была получена за один год, а 250%-ная доходность — более чем за 10 лет, то первая лучше. В то же время, если первая доходность ассоциируется с 10%-ным риском в годовом исчислении, а вторая — с 50%-ным риском, то системы эквивалентны. Соотношение доходности и риска крайне важно для результативности, о чем мы поговорим в главе 21 «Тестирование систем». А пока важно лишь запомнить, что доходность и риск следует выражать в годовом исчислении или нормализовать иным образом, чтобы сравнения имели смысл.

Расчет доходности

Расчет как однопериодной, так и годовой доходности является неотъемлемой частью всех оценок результативности. В простейшей форме *однопериодная доходность* R , или *доходность за период владения*, часто представляется как

$$\begin{aligned} \text{Доходность} &= \frac{\text{конечное значение} - \text{начальное значение}}{\text{начальное значение}} \\ &= \frac{\text{конечное значение}}{\text{начальное значение}} - 1. \end{aligned}$$

Для фондового рынка, где цены меняются непрерывно, это можно записать как

$$r_1 = \frac{p_1 - p_0}{p_0} = \frac{p_1}{p_0} - 1,$$

где p_0 — первоначальная цена, а p_1 — цена после истечения одного периода.

В индустрии ценных бумаг нередко предпочитают другой расчет:

$$r_n = \ln \left(\frac{P_t}{P_{t-1}} \right).$$

Оба метода имеют свои достоинства и недостатки. Ни один не относится к категории «правильных» расчетов. Отметим, что в некоторых компьютерных программах функция $\log$ фактически является натуральным логарифмом, а $\log 10$ — десятичным логарифмом. Всегда лучше перепроверить определения. Чтобы различать эти два расчета, первый метод будем называть *стандартным методом*, а второй — *логарифмическим методом*.

В примере, показанном в табл. 2.3 и охватывающем 22 дня, стандартная доходность находится в столбце D, а логарифмическая — в столбце E. Различия кажутся небольшими, но средние имеют значения 0,00350 и 0,00339. Стандартная доходность оказывается больше на 3,3% всего за один месяц торговли. С такими темпами за год стандартный метод принесет доходность на 40% больше. *Чистая стоимость активов* (net asset value — NAV), активно используемая в этой книге, определяется с учетом доходности за период и чаще всего начинается со значения $NAV_0 = 100$.

$$NAV_t = NAV_{t-1} \times (1 + r_t).$$

Доходность в годовом исчислении

В большинстве случаев лучше всего нормализовать доходность путем ее приведения к *годовому исчислению*. Это особенно полезно при сравнении результатов двух тестов, когда каждый охватывает свой период времени. При приведении к годовому исчислению важно знать следующее.

- К правительственным инструментам применяется годовая база, равная 360 дням (на основе 90-дневных кварталов).
- Для большинства других данных, которые могут изменяться ежедневно, характерна годовая база 365 дней.
- Торговую доходность лучше рассчитывать на основе 252-дневной годовой базы, это типичная продолжительность торгового года в США (262 дня в Европе).

В следующих формулах используется годовая база 252 дня, и это будет стандартом для данной книги, за исключением расчетов некоторых процентных ставок, однако в формулы всегда можно подставить 365, или 360, или даже 260 торговых дней, в зависимости от условий в других частях света.

Доходность в годовом исчислении (annualized rate of return — AROR) на основе простого процента для инвестиции в течение n дней равна

$$AROR_{\text{простая}} = \frac{E_n}{E_0} \times \frac{252}{n},$$

где E_0 — начальный капитал или остаток на счете, E_n — капитал в конце периода n , а $252/n$ — годовая база в виде десятичной дроби. Когда однопериодная доходность рассчитывается с использованием стандартного метода, доходность в годовом исчислении на основе сложного процента равна

$$AROR_{\text{сложная}} = \left(\frac{E_n}{E_0} \right)^{\frac{252}{n}} - 1.$$

Таблица 2.3. Расчет доходности и чистой стоимости активов на основе дневных прибылей и убытков

Дата	П/У	Кумулятивная П/У	Стандартная доходность	Логарифми- ческая доходность	NAV	Логарифми- ческая NAV
10.9.2010		100 000			100,00	100,00
13.9.2010	1154	101 154	0,01154	0,01147	101,15	101,15
14.9.2010	1795	102 949	0,01774	0,01759	102,95	102,91
15.9.2010	-1859	101 090	-0,01806	-0,01822	101,09	101,08
16.9.2010	-1603	99 487	-0,01585	-0,01598	99,49	99,49
17.9.2010	449	99 936	0,00451	0,00450	99,94	99,94
20.9.2010	1090	101 026	0,01090	0,01084	101,03	101,02
21.9.2010	2949	103 974	0,02919	0,02877	103,97	103,90
22.9.2010	1346	105 320	0,01295	0,01286	105,32	105,18
23.9.2010	64	105 384	0,00061	0,00061	105,38	105,24
24.9.2010	-2051	103 333	-0,01946	-0,01966	103,33	103,28
27.9.2010	3269	106 602	0,03164	0,03115	106,60	106,39
28.9.2010	1795	108 397	0,01684	0,01670	108,40	108,06
29.9.2010	-1154	107 243	-0,01064	-0,01070	107,24	106,99
30.9.2010	128	107 372	0,00120	0,00119	107,37	107,11
01.10.2010	-705	106 666	-0,00657	-0,00659	106,67	106,45
04.10.2010	1090	107 756	0,01022	0,01016	107,76	107,47
05.10.2010	-449	107 307	-0,00416	-0,00417	107,31	107,05
06.10.2010	2308	109 615	0,02150	0,02128	109,62	109,18
07.10.2010	-769	108 846	-0,00702	-0,00704	108,85	108,48
08.10.2010	-256	108 589	-0,00236	-0,00236	108,59	108,24
12.10.2010	-1218	107 372	-0,01122	-0,01128	107,37	107,11
Среднее			0,00350	0,00339		
Стд. откл.			0,01502	0,01495		
Год. стд. откл.			0,23846	0,23730		
AROR					125,85%	119,69%

Обратите внимание, что $AROR$, или R (капитал), является значением в годовом исчислении, а r — это дневная или однопериодная доходность. Кроме того, форма представления результатов в этих двух расчетах разная. В случае простой доходности увеличение на 25% выглядит как 0,25, а в случае сложной доходности такой же прирост записывается как 1,25.

Если к однопериодной доходности применить логарифмический метод, то доходность в годовом исчислении будет равна сумме доходностей, деленной на количество лет:

$$AROR_{\text{логарифмическая}} = \frac{\sum_{i=1}^n r_i}{n}.$$

Пример этого можно найти в строке $AROR$ столбца F предыдущей таблицы. Обратите внимание на то, что при использовании логарифмического метода доходность в годовом исчислении оказывается намного ниже, чем при использовании деления и сложных процентов. В настоящей книге везде используется сложная доходность.

Вероятность доходности

Стандартное отклонение и доходность в сложных процентах используются совместно для определения вероятности достижения целевой доходности. В следующем расчете¹ среднее арифметическое непрерывной доходности равно $\ln(1 + R_g)$, при этом предполагается, что доходность имеет нормальное распределение.

$$z = \frac{\ln\left(\frac{T}{B}\right) - \ln(1 + R_g)n}{\sqrt{s \times n}},$$

где z — нормализованная переменная (см. приложение A1);

T — целевое значение, или целевая доходность;

B — начальная стоимость инвестиции;

R_g — среднее геометрическое периодической доходности;

n — количество периодов;

s — стандартное отклонение логарифмов количеств 1 плюс периодическая доходность.

Риск и волатильность

Хотя думать о доходности всегда приятно, не следует забывать об оценке риска. В связи с этим следует упомянуть о двух экстремальных рисках. Первым из них является *событийный риск*, который принимает форму непредсказуемого скачка цен. Худший его вариант — *катастрофический риск*, приносящий непоправимый ущерб или крах. Второй риск связан с *чрезмерным использова-*

¹ Это и другие очень ясные объяснения доходности см.: Peter L. Bernstein, *The Portable MBA in Investment*.

нием заемных средств, или *левериджа* в портфеле, что приводит к краху в случае череды неудачных сделок. Риски, связанные со скачками цен и *левериджем*, подробно описываются в последующих главах.

Измерение нормального риска полезно для сравнения результативности двух систем. Обычно оно применяется при сравнении доходности отдельной акции или всего портфеля с ориентиром, таким как доходность S&P 500 или какого-нибудь облигационного фонда. Самым распространенным методом оценки риска является стандартное отклонение (σ) доходности (r), о чем рассказывалось в этой главе ранее. В большинстве случаев при описании риска стандартное отклонение называют также *волатильностью*. Когда мы говорим о *целевой волатильности* портфеля, то имеем в виду процент риска, представленный одним стандартным отклонением доходности в годовом исчислении. Например, в предыдущей таблице столбцы D и E показывают дневные доходности. Стандартные отклонения этих доходностей показаны в тех же столбцах в строке «Стд. откл.» — это значения 0,01512 и 0,01495. Если ограничиться только столбцом D, то одно стандартное отклонение 0,01502 означает, что существует 68%-ная вероятность получения дневной доходности или убыточности в размере менее 1,502%. Однако целевая волатильность всегда соотносится с годовым риском, и, чтобы преобразовать дневную доходность в годовую, мы просто умножим ее на $\sqrt{252}$. Тогда дневное стандартное отклонение доходности в размере 1,512% становится *годовой волатильностью* 23,8%, которая также показана в нижней части нашей таблицы. Поскольку нас волнует только риск убытка, существует 16%-ная вероятность, что в течение года мы можем потерять 23,8%. Чем больше стандартное отклонение доходности, тем больше риск.

Бета

Бета (β) обычно используется в индустрии ценных бумаг для выражения взаимосвязи отдельного рынка с индексом или портфелем. Если бета равна нулю, никакой взаимосвязи нет. Если она положительна, значит, отдельный ряд движется вместе с индексом, или выше, или ниже него. Когда бета увеличивается, волатильность отдельного рынка становится все больше, чем у индекса. А именно:

$0 < \beta < 1$ — волатильность отдельного рынка меньше, чем у индекса;

$\beta = 1$ — волатильность отдельного рынка такая же, как у индекса;

$\beta > 1$ — волатильность отдельного рынка больше, чем у индекса.

Отрицательная бета подобна отрицательной корреляции, когда движения рынка противоположны движению индекса.

Чтобы найти бету, надо рассчитать линейную регрессию отдельного рынка по индексу. Она равна наклону отдельного рынка, деленному на наклон индекса. Дополнительная величина *альфа* является свободным членом решения. Результат можно получить, используя Excel, подробное описание см. в главе 6. Общая формула беты выглядит так:

$$\text{Бета}(A) = \frac{\text{cov}(\text{доходность } A, \text{доходность } B)}{\text{var}(\text{доходность } B)},$$

где A — отдельный рынок, а B — портфель или индекс.

Корректировка для возврата к целевой волатильности

Если мы имеем целевую волатильность 12%, т. е. готовы принять 16%-ную вероятность потери через год 12%, но фактическая доходность показывает для инвестиции в \$ 100 000 годовую волатильность 23,8%, то для возврата к 12%-ной цели мы просто увеличиваем инвестицию в $23,8/12,0$ раз, или на коэффициент 1,98 до уровня \$ 198 000, сохраняя прежний размер позиции. По существу, вы уменьшаете леверидж своих позиций, ограничивая торговлю меньшим процентом величины своего счета. Или же можно сократить размер позиции, разделив ее на 1,98, и сохранить размер инвестиции без изменения.

Все результаты в этой книге демонстрируются при целевой волатильности 12%, если не указано иное. Это считается скромным уровнем риска, который у некоторых хедж-фондов, например, может достигать 18%. Это позволяет сравнивать различные системы и результаты тестирования и видеть их на уровне риска, который с наибольшей вероятностью представляет целевые результаты торговли.

Представление дневной и месячной доходности в годовом исчислении

В предыдущих примерах использовались дневные данные и коэффициент $\sqrt{252}$. Для месячных данных, которые встречаются в публикуемых таблицах результативности чаще всего, следует брать месячную доходность и умножать ее на $\sqrt{12}$. В принципе, преобразование в годовое исчисление можно выполнять, умножая данные на квадратный корень из количества элементов данных в году. Поэтому мы используем 252 для дневных данных, 12 для месячных, 4 для квартальных и т. д.

Месячные данные всегда кажутся менее волатильными

Как правило, в финансовых отчетах используются месячные данные. Нужно помнить, что такой подход работает в пользу тех, кто публикует результаты. Маловероятно, чтобы дневная *чистая стоимость актива* (NAV) достигла своего максимального или минимального уровня в последний день месяца, поэтому экстремальные значения редко оказываются в поле зрения, и статистика результатов кажется более сглаженной, чем при использовании дневной доходности. Поэтому тем, кто выполняет процедуру «дью дилидженс» перед инвестированием, нередко нужны дневные данные, чтобы не пропустить большую просадку в середине месяца.

Воспользовавшись в качестве примера индексом S&P, мы видим, что в 1990–2010 гг. годовая волатильность дневной доходности составляла 18,6%, а месячной доходности — только 15,3%. Риск, рассчитанный на основе месячной доходности, на 17,7% ниже, однако доходность в годовом исчислении оказывается одинаковой, потому что в ней используются только начальные и конечные значения.

Риск убытка

Поскольку стандартное отклонение симметрично, появление любого ряда скачков в прибыли интерпретируется как увеличение риска. Некоторые аналитики

полагают, что правильнее измерять риск, ограничиваясь только случаями снижения доходности, или просадки. Использование одних только убытков называют *нижними частичными моментами*, где *нижние* означает риск снижения доходности, а *частичные* означает, что используется только одна сторона распределения доходности. Легче всего увидеть это на примере *полудисперсии*, которая измеряет дисперсию, находящуюся ниже среднего или некоторого целевого значения,

$$\text{Полудисперсия} = \frac{\sum_{i=1}^n (\bar{R} - r_i)^2}{n}, \text{ где каждая } r_i < \bar{R}.$$

Однако чаще всего для определения результативности системы используются дневные просадки, т. е. чистый убыток в каждый день, когда величина капитала оказывается ниже его пикового значения. Например, если в день t капитал системы вырос до \$ 25 000, а затем был получен дневной убыток \$ 500, за которым последовал еще один убыток \$ 250, у нас будет два входных значения, $\frac{500}{25000}$ и $\frac{750}{25000}$, или 0,02 и 0,03. В расчете полудисперсии используются только те чистые прибыли, которые ниже недавних пиков. Альтернативно можно просто взять стандартное отклонение этих дневных просадок, чтобы найти их вероятную величину.

Использование одних лишь просадок для предсказания других просадок влечет за собой проблему, которая состоит в том, что такой расчет ограничивает количество рассматриваемых случаев и отбрасывает вероятность влияния прибыли выше нормальной на повышение общего уровня риска. В ситуациях, когда количество данных в тестах ограничено, использование совокупности прибылей и убытков дает более надежные результаты. Но, когда данных очень много, использование одних только просадок может приносить очень хорошие плоды.

Полное описание методов измерения результатов можно найти в главе 21 «Тестирование систем», а также в разделах главы 23 «Измерение доходности и риска» и «Индекс язвы».

Индекс

Индексы предназначены для того, чтобы обобщать индивидуальные особенности. При этом данные часто сглаживаются, и из них извлекается полезная информация. В последние годы индексы завоевали огромную популярность. Если в начале 1980-х гг. на фьючерсных рынках торговались лишь Value Line и S&P 500, то теперь существуют фьючерсные контракты на фондовые индексы, представляющие рынки любой промышленно развитой страны. Создание фондов, таких как SPDR (прозванные «спайдерами», на основе S&P 500), Diamonds (DIA, на основе промышленных индексов Dow Jones) и Qs (QQQ, на основе NASDAQ 100), дало трейдерам знакомый инструмент для инвестирования в широкий рынок вместо отдельных акций. Отраслевые секторы, такие как фармацевтика, здравоохранение и технологии, которые были представлены сначала взаимными фондами, затем ETF, теперь могут торго-

ваться и на фьючерсном рынке. Все эти индексные рынки имеют дополнительное преимущество, заключающееся в том, что на них нет ограничений в виде необходимости заимствовать акции для короткой продажи или правила «плюс тик» (если его восстановят), требующего совершения коротких продаж только на росте цены.

Индексные рынки позволяют как индивидуальным, так и институциональным участникам осуществлять ряд специализированных инвестиционных стратегий. Они могут покупать и продавать широкий рынок, могут переключаться с одного сектора на другой (ротация секторов) или продавать переоцененный сектор, покупая широкий рыночный индекс (статистический арбитраж). Институты находят очень желательным с точки зрения и затрат, и налогов временно хеджировать свои портфели акций, продавая фьючерсы на S&P 500 вместо того, чтобы ликвидировать позиции в акциях. Они могут также хеджироваться, используя опционы на фьючерсы S&P или SPY. Индекс упрощает процесс принятия решений в торговле. Если индекс не существует, его можно построить самостоятельно, в зависимости от поставленных целей.

Индекс играет важную роль как ориентир результативности. Большинство инвесторов полагает, что торговая программа привлекательна только в том случае, если ее соотношение доходности и риска лучше, чем у портфеля, состоящего на 60% из акций (представленных индексом S&P 500) и 40% из облигаций (Lehman Brothers Treasury Index). Если ваши результаты лучше, чем у индекса, значит, вы *создали альфу* и превзошли рынок.

Построение индекса

Индекс представляет собой стандартизированный способ выражения движения цены, обычно в виде накопления процентных изменений. Большинство индексов имеет стартовое значение 100, зафиксированное на определенную дату. Выбор базисного года нередко «произволен», но может приходиться и на период стабильности цен. В США базисным годом для производительности труда и безработицы является 1982 г., для уверенности потребителей — 1985 г., а для совокупности опережающих индикаторов — 1987 г. Ежегодник CRB Yearbook показывает индекс промышленных цен (PPI) начиная с 1913 г. Например, PPI, публикуемый ежемесячно, имел в октябре 2010 г. значение 186,8, а в сентябре 2010 г. — 185,1, продемонстрировав прирост на 0,9184% за месяц. Величина индекса менее 100 означает, что индекс стал меньше, чем тогда, когда началось его исчисление.

Каждое значение индекса рассчитывается на основе предыдущего значения следующим образом:

Текущее значение индекса = Предыдущее значение индекса ×

$$\times \left(\frac{\text{Текущая цена}}{\text{Предыдущая цена}} \right),$$

а однопериодная доходность рассчитывается так же, как было показано в этой главе ранее.