

ПРОГНОЗНЫЕ МОДЕЛИ СУБЪЕКТИВНЫХ ПРЕДПОЧТЕНИЙ

В. В. Давнис, В. И. Тинякова

Воронежский государственный университет

ВВЕДЕНИЕ

В настоящее время все чаще при обосновании управленческих решений в качестве альтернативы условиям неопределенности используются различные схемы субъективных предпочтений. Однако, несмотря на актуальность данного подхода и разнообразие его схем, в практике обоснования по преимуществу используются методы из одного и того же давно известного набора, основу которого составляет простое ранжирование, балльное оценивание и парные сравнения. Всем, кто знаком с процедурами экспертного оценивания, известно, что длительное время этот набор не пополняется новыми методами. И это, на наш взгляд, не случайно.

Главная особенность данных методов в том, что для числового представления получаемых результатов, в основном, используются номинальные и ранговые шкалы. В этих шкалах и кроется тот консерватизм, который утвердился в отношении методов обработки экспертной информации. Его природа очевидна: низкая разрешающая способность экспертов, служащая непреодолимым барьером для повышения точности экспертных оценок. Возникает естественный вопрос: «Какой смысл в разработке новых подходов и более точных методов, если они из-за указанного барьера не приводят к уточнению финальных результатов?» Трудно возразить этому тезису. И все же смысл есть. Он появляется в тех случаях, когда меняется привычное представление о сути решаемых задач.

1. ЭКОНОМЕТРИЧЕСКИЕ РЕШЕНИЯ В ЭКСПЕРТНОМ ОЦЕНИВАНИИ

Новый взгляд на проблему экспертного оценивания является той отправной точкой,

в которой начинаются исследования по формированию результатов экспертного опроса в соответствии с этим взглядом. Основная идея подобных исследований в том, чтобы обосновать необходимость и полезность получения результатов опроса не в ранговых и номинальных шкалах, как это принято, а в виде моделей, концентрирующих в своей аналитической структуре субъективные предпочтения экспертов. В качестве примера, демонстрирующего новый взгляд на обработку экспертной информации, рассмотрим, ставшую классической, задачу ранжирования показателей по их степени влияния на возможность появления какого-либо события.

Задача ранжирования является одной из наиболее распространенных в практике экспертного оценивания. Ее решение можно получить с помощью любого из вышеуказанных методов. Однако, несмотря на многообразие методов, суть используемого в них подхода одна — непосредственное оценивание показателей. Наряду с простотой реализации этот подход имеет и ряд недостатков. Очевидно, что его применение имеет смысл только в линейном случае, когда степень влияния не зависит от структуры оцениваемого набора показателей.

В реальных ситуациях все гораздо сложнее. Представление о линейном взаимодействии скорее абстракция, помогающая упростить задачу, сделав ее всегда решаемой, но с некоторой ошибкой, которой можно пренебречь. Логика получения результатов по такой схеме оценивания без учета совместных эффектов вполне объяснима. Решение ищется для конкретной ситуации с фиксированной структурой показателей, которая хотя и не указывается в задании эксперту, но, как правило, присутствует в его представлениях о решаемой задаче. Но как

только структура начинает изменяться, сразу же появляются неучтенные эффекты взаимодействия и надежность экспертных оценок резко снижается. Поэтому схему непосредственного оценивания показателей необходимо заменить другой, основанной на модельном представлении результатов в зависимости от структуры, но без усложнения самой процедуры опроса экспертов. При этом модель, отражающая взаимосвязь между возможностью появления интересующего нас события и набором оцениваемых показателей, должна быть по всей вероятности нелинейной и, кроме того, эконометрической, так как интерес вызывает не только механизм взаимодействия, но и адекватное отражение количественной оценки силы этого взаимодействия. Полученные оценки в виде коэффициентов регрессии, как раз и содержат информацию о ранговой структуре показателей. Кроме отмеченного, в подобном подходе реализуется естественное желание заменить повторные экспертные опросы прогнозными оценками. Последнее является важной особенностью. Именно этой возможностью не обладают прямые методы экспертного оценивания.

Таким образом, смысл рассматриваемого здесь подхода в том, чтобы экспертную информацию использовать для построения модели, с помощью которой будут получены интересующие нас оценки, а не для непосредственного получения самих оценок. Возникает естественный вопрос, каким образом экспертная информация может использоваться в этих целях. По всей видимости, можно предложить несколько подходов, обеспечивающих реализацию обсуждаемой здесь идеи. Наше предложение заключается в том, чтобы интуицию и знания экспертов применить для формирования специальных выборочных совокупностей, которые мы будем называть псевдовыборками. Данные сформированных псевдовыборочных совокупностей используются для оценивания коэффициентов регрессионной модели, представляющей собой инструмент многопланового применения: анализ, оценка значимости факторов, прогноз ожидаемых событий и т.п. Естественно, это значительно расширяет область применения экспертных решений.

Формальная реализация данного подхода предполагает введение бинарной переменной со следующим смыслом:

$$y = \begin{cases} 1, & \text{если по мнению эксперта событие} \\ & \text{должно произойти,} \\ 0, & \text{в противном случае.} \end{cases}$$

Будем считать, что значение этой переменной, характеризующей появление интересующего нас события, зависит от оцениваемого нами набора показателей $x_1, x_2, \dots, x_m$ и существует некоторое множество различных вариантов $x_1, x_2, \dots, x_n$ этих наборов $x_i = (x_{i1}, x_{i2}, \dots, x_{im})$, отличающихся друг от друга всеми или некоторыми своими компонентами (оцениваемыми показателями). Предполагается, что у каждого эксперта есть представление о том, при реализации каких вариантов ожидаемое событие будет иметь место, а при реализации каких нет. Математически это предположение записывается в виде зависимости

$$y_i^k = f(x_{i1}, x_{i2}, \dots, x_{im}) + \varepsilon_i^k, \quad (1)$$

где y_i^k — ожидаемое значение бинарной зависимой переменной, которое k -ый эксперт связывает с i -ым набором оцениваемых показателей; $f(x_{i1}, x_{i2}, \dots, x_{im})$ — индексная функция, т.е. функция, принимающая всего два значения 0 и 1; ε_i^k — ошибка, которую может допустить k -ый эксперт, оценивая влияние i -го набора на появление ожидаемого события.

Теперь становится понятной реализация основанной на модельном подходе идеи получения экспертных решений. Сначала в результате целевого опроса экспертов формируется псевдовыборка, объединяющая в себе субъективные мнения по поводу интересующих нас закономерностей, предпочтений, рейтингов, прогнозных оценок и т.п. Затем по данным псевдовыборки строится регрессионная зависимость (1), связывающая субъективные мнения с одновременным их усреднением в единую формализованную зависимость. Построенная таким образом модель, по сути, является концентрированным выражением обобщенного мнения экспертов по изучаемой проблеме и может использоваться для анализа и получения всевозможных оценок.

Практическая реализация этой процедуры требует рассмотрения целого ряда довольно сложных вопросов:

1) Как построить регрессию на бинарную переменную?

2) Какими способами эксперты должны формировать выборочную совокупность (псевдовыборку) для построения регрессии с бинарной зависимой переменной?

3) Как оценить компетентность эксперта и адекватность построенной таким образом модели?

4) Каким образом проверить согласованность мнений опрашиваемой группы экспертов?

5) Как оценить надежность расчетных характеристик, получаемых с помощью регрессии субъективных предпочтений?

6) Какую содержательную интерпретацию имеют оценки, полученные в результате моделирования экспертных предпочтений?

7) Можно ли, а если можно, то как использовать статистические методы для проверки различного рода гипотез, выдвигаемых относительно оцениваемых параметров и объектов?

По сути, в этих вопросах нет ничего нового. Все они в той или иной степени присутствуют в задачах прямого экспертного оценивания, обеспечивая надежность получаемых результатов. Но расчет и анализ характеристик, а также соответствующие выводы на их основе в новом подходе отличаются от того, как это делается в традиционном, и поэтому требуют специального рассмотрения.

2. МОДЕЛЬ БИНАРНОГО ВЫБОРА И МЕТОДЫ ЕЕ ПОСТРОЕНИЯ

Реализация эконометрического подхода требует уточнения вида индексной функции (1). Учитывая область значений индексной функции и вероятностную природу экспертных суждений, представим ее следующим образом:

$$y_i = \begin{cases} 1, & \text{если } F(\mathbf{x}_i, \mathbf{b}) > 0.5, \\ 0, & \text{если } F(\mathbf{x}_i, \mathbf{b}) \leq 0.5, \end{cases} \quad (2)$$

где $F(\mathbf{x}_i, \mathbf{b})$ — функция распределения вероятности того, что при данном комплексе условий $\mathbf{x}_i$ событие произойдет, т.е. $y_i = 1$.

Тогда эконометрическую модель можно будет представить в виде зависимости между бинарной переменной, отражающей факт появления события, и предсказанной на основе суждений экспертов вероятности его реализации

$$y_i = F(\mathbf{x}_i, \mathbf{b}) + \varepsilon_i. \quad (3)$$

В практике решения прикладных задач в качестве F чаще всего используются две функции. Если такой функцией является стандартная нормальная вероятностная функция распределения

$$\Phi(\mathbf{x}_i, \mathbf{b}) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\mathbf{x}_i, \mathbf{b}} e^{-\frac{z^2}{2}} dz, \quad (4)$$

то регрессионная зависимость называется пробит-моделью. В случае, когда используется логистическая функция

$$\Lambda(\mathbf{x}_i, \mathbf{b}) = \frac{1}{1 + e^{-\mathbf{x}_i, \mathbf{b}}}, \quad (5)$$

зависимость называется логит-моделью.

Построение регрессионных моделей с использованием нелинейных зависимостей подобного типа практически исключает применение метода наименьших квадратов, поэтому рассмотрим этот вопрос подробнее. Для оценивания моделей бинарного выбора удобно использовать метод максимального правдоподобия. Применение этого метода осуществляется в предположении, что каждое наблюдение может трактоваться как однократный выбор из распределения Бернулли. Модель с вероятностью успеха $F(\mathbf{x}_i, \mathbf{b})$ и независимыми наблюдениями (эксперты опрашиваются независимо друг от друга) представляет собой вероятность совместного появления всей совокупности ожидаемых событий

$$\begin{aligned} P(Y_1 = y_1, Y_2 = y_2, \dots, Y_n = y_n) &= \\ &= \prod_{y_i=1} F(\mathbf{x}_i, \mathbf{b}) \prod_{y_i=0} (1 - F(\mathbf{x}_i, \mathbf{b})). \end{aligned} \quad (6)$$

Для каждого вектора $\mathbf{y}$, представляющего собой результаты конкретного экспертного опроса, величина вероятности зависит от вектора оцениваемых параметров $\mathbf{b}$ и может быть записана как функция правдоподобия

$$L(\mathbf{y}, \mathbf{b}) = \prod_{i=1}^n F(\mathbf{x}_i, \mathbf{b})^{y_i} [1 - F(\mathbf{x}_i, \mathbf{b})]^{1-y_i}. \quad (7)$$

В данной форме записи множители произведения селектируются с помощью компонент вектора $\mathbf{y}$, принимающих всего два значения 0 или 1.

Удобнее и математически проще максимизировать логарифмическую функцию правдоподобия

$$\ln L = \sum_{i=1}^n [y_i \ln F(\mathbf{x}_i, \mathbf{b}) + (1 - y_i) \ln(1 - F(\mathbf{x}_i, \mathbf{b}))]. \quad (8)$$

Используя сокращенные записи $F_i = F(\mathbf{x}_i, \mathbf{b})$ и $F'_i(\mathbf{x}_i, \mathbf{b}) = f_i$, выпишем для логарифмической функции правдоподобия условия максимизации первого порядка

$$\frac{\partial \ln L}{\partial \mathbf{b}} = \sum_{i=1}^n \left[\frac{y_i f_i}{F_i} + (1 - y_i) \frac{-f_i}{(1 - F_i)} \right] \mathbf{x}'_i = 0. \quad (9)$$

Подставляя в (9) логистическое и нормальное распределения, получаем системы нелинейных уравнений соответственно для оценки коэффициентов логит-модели

$$\sum_{i=1}^n (y_i - \Lambda_i) \mathbf{x}_i = 0 \quad (10)$$

и пробит-модели

$$\sum_{i=1}^n \left[\frac{q_i \phi(q_i \mathbf{x}_i, \mathbf{b})}{\Phi(q_i \mathbf{x}_i, \mathbf{b})} \right] \mathbf{x}'_i = 0. \quad (11)$$

В (11) использовано обозначение $q_i = 2y_i - 1$.

Нелинейность полученных систем требует применения численных методов. Прежде чем приступить к численному решению покажем, что логарифмическая функция правдоподобия строго вогнута и, следовательно, имеет единственный максимум.

Основным признаком строгой вогнутости является отрицательность второй производной. Для логистической функции имеет место очевидное неравенство

$$\frac{d^2(\ln F(x))}{dx^2} = \frac{-e^{-x}}{(1 + e^{-x})^2} < 0. \quad (12)$$

В случае нормального распределения получается тот же самый результат. Таким образом, решение системы (10) или (11) приводит к получению оценок, максимизирующих соответствующие функции правдоподобия, т.е. другими словами, существует метод, обеспечивающий построение модели субъективных предпочтений в такой постановке.

3. ЧИСЛЕННОЕ РЕШЕНИЕ УРАВНЕНИЙ ПРАВДОПОДОБИЯ

Для решения уравнения правдоподобия

$$\frac{\partial \ln L(\mathbf{b})}{\partial \mathbf{b}} = 0 \quad (13)$$

удобно использовать метод Ньютона—Рафсона. Описание вычислительной схемы проведем для общего случая без уточнения, на основе какого распределения была построена функция правдоподобия.

Считая левую часть системы (13) дифференцируемой вектор-функцией (для исследуемых здесь распределений это действительно так), запишем отрезок ряда Тейлора, являющегося линейной аппроксимацией этой функции в окрестности некоторой точки $\mathbf{b}_0$

$$\frac{\partial \ln L(\mathbf{b})}{\partial \mathbf{b}} = \frac{\partial \ln L(\mathbf{b}_0)}{\partial \mathbf{b}} + \frac{\partial^2 \ln L(\mathbf{b}_0)}{\partial \mathbf{b} \partial \mathbf{b}'} (\mathbf{b} - \mathbf{b}_0). \quad (14)$$

Точка $\mathbf{b}_0$ является начальным приближением искомой оценки и ее значение можно определить как вектор параметров линейной регрессии с помощью метода наименьших квадратов.

Обозначив произвольную точку окрестности через $\mathbf{b}_1$ и помня, что нашей целью является нахождение такого вектора параметров, который обращает первую производную в ноль, запишем

$$\mathbf{0} = \frac{\partial \ln L(\mathbf{b}_0)}{\partial \mathbf{b}} + \frac{\partial^2 \ln L(\mathbf{b}_0)}{\partial \mathbf{b} \partial \mathbf{b}'} (\mathbf{b}_1 - \mathbf{b}_0). \quad (15)$$

Раскрывая круглые скобки и перенося влево член, содержащий $\mathbf{b}_1$, а затем, умножая обе части уравнения на обратную матрицу, получаем итерационный процесс нахождения искомого решения

$$\mathbf{b}_{k+1} = \mathbf{b}_k - \left\{ \frac{\partial^2 \ln L(\mathbf{b}_k)}{\partial \mathbf{b} \partial \mathbf{b}'} \right\}^{-1} \frac{\partial \ln L(\mathbf{b}_k)}{\partial \mathbf{b}}. \quad (16)$$

Последовательность $\{\mathbf{b}_k\}$ сходящаяся, и можно показать, что ее предел является решением системы (13).

4. СТАТИСТИЧЕСКАЯ ЗНАЧИМОСТЬ КОЭФФИЦИЕНТОВ

Проверка статистической значимости коэффициентов модели осуществляется с помощью статистики Вальда, построенной с помощью стандартных ошибок коэффи-

циентов модели, в качестве которых используются корни квадратные из диагональных элементов ковариационной матрицы оценок $\hat{\mathbf{b}}$. Ковариационная матрица оценок модели бинарного выбора равняется матрице, обратной к информационной матрице Фишера, которая представляет собой математическое ожидание Гесса, взятое с обратным знаком

$$V(\hat{\mathbf{b}}) = I^{-1} = \left\{ -E \left[\frac{\partial^2 \ln(L)}{\partial \mathbf{b} \partial \mathbf{b}'} \right] \right\}^{-1}. \quad (17)$$

Предполагается, что математическое ожидание здесь берется условно по $\mathbf{x}$.

Как нетрудно понять, асимптотическая матрица ковариации зависит от неизвестного параметра $\mathbf{b}$. Поэтому непосредственное использование ее в практических расчетах исключено. Рекомендации здесь те же самые, что и при использовании обычной регрессии: неизвестные параметры, присутствующие в матрице следует заменить соответствующими оценками. Руководствуясь этим общим правилом, можно записать

$$V(\hat{\mathbf{b}}) = \left\{ -E \left[\frac{\partial^2 \ln(L)}{\partial \mathbf{b} \partial \mathbf{b}'} \right] \right\}_{\mathbf{b}=\hat{\mathbf{b}}}^{-1} \quad (18)$$

и использовать в практических расчетах не ковариационную матрицу, а ее оценку:

$$\hat{V}(\hat{\mathbf{b}}) = I^{-1} = \left\{ \sum_{i=1}^n \frac{[f(\mathbf{x}_i, \hat{\mathbf{b}})]^2}{F(\mathbf{x}_i, \hat{\mathbf{b}})[1 - F(\mathbf{x}_i, \hat{\mathbf{b}})]} \mathbf{x}_i' \mathbf{x}_i \right\}^{-1}. \quad (19)$$

Если через S_{kk} обозначить стандартные ошибки соответствующих оценок $\hat{\mathbf{b}}_k$, то статистика Вальда записывается в виде

$$w_k = \left(\hat{\mathbf{b}}_k / S_{kk} \right)^2, \quad k = 0, 1, \dots, m. \quad (20)$$

Сравнение ее с табличным значением $\chi_{95\%}^2(n-1)$ позволяет установить значимость коэффициентов модели бинарного выбора.

5. ПРИНЦИПЫ ФОРМИРОВАНИЯ ПСЕВДОВЫБОРОЧНЫХ СОВОКУПНОСТЕЙ

Вопрос применения для наших целей моделей бинарного выбора имеет как бы два аспекта. Первый, связанный с методами построения этих моделей, со всеми деталями был рассмотрен выше, и не является ак-

туальным. Единственное требование — корректность применения. Второй аспект требует обоснования принципов формирования псевдовыборочных совокупностей.

Смысл основной задачи, стоящей перед экспертами, формирующими псевдовыборку, в том, чтобы на данные выборочного множества перенести собственные представления о механизмах предполагаемых закономерностей между объясняющими переменными и ожидаемыми событиями. Тогда, если проведение такой процедуры было успешным, то по замыслу построенная модель должна отражать ту закономерность, руководствуясь которой эксперт оценивал степень воздействия выборочных значений на возможные проявления интересующего нас события. *Таким образом, главное отличие псевдовыборки от выборки в том, что в ее данных содержится информация, которую эксперты сумели обнаружить и связать своими субъективными оценками со значениями зависимой переменной.*

Способы формирования псевдовыборки зависят от смыслового содержания решаемой задачи. Это не совсем строго формализованные процедуры и поэтому для их успешного применения в каждом конкретном случае требуется адаптивное вмешательство. Рассмотрим типовые ситуации, которые возникают в практике анализа и прогнозирования бизнес процессов. Можно выделить три ситуации, которые отличаются принципами формирования псевдовыборки.

В первой ситуации псевдовыборку формируют непосредственно из выборочной совокупности с известными значениями зависимой (дихотомической) и независимыми переменными. Это тот случай, когда взаимосвязь между y_i и наборами $\mathbf{x}_i$ ($i = 1, n$) существует, но эксперты сомневаются в абсолютной правомерности и надежности такой связи. Применяя принцип усиления взаимосвязей субъективными мнениями, они своими оценками уточняют возможность появления соответствующего значения y_i при заданном наборе $\mathbf{x}_i$. Уточняющие оценки удобно, хотя и не обязательно, получать в баллах. Важно только, чтобы результаты экспертного опроса можно было интерпретировать как частоту (число случаев), с которой данное наблюдение тиражируется в выборочной совокупности.

Вторая ситуация, которую мы намерены рассмотреть, предусматривает случаи, когда выборочное множество состоит только из объясняющих переменных, и требуется на основе принципа *субъективной идентификации взаимосвязей* восстановить значения бинарной зависимой переменной. В соответствии с этим принципом эксперт каждому набору объясняющих переменных выборочной совокупности ставит в соответствие одно из возможных значений 0 или 1 зависимой переменной.

В сформированной таким образом псевдовыборке, как нетрудно понять, содержится информация, отражающая субъективное мнение по поводу того, какие условия, описываемые объясняющими переменными, благоприятствуют появлению события, а какие — нет. Другими словами, экспертами сгенерированы значения бинарной переменной в зависимости от значений объясняющих переменных. Естественно, модель, построенная по данным так сформированной псевдовыборки, будет являться довольно грубым приближением к той зависимости, которую эксперты пытались описать своими предпочтениями в процессе формирования дискретной зависимой переменной. Поэтому желательно провести некоторые уточнения, используя для этого, например, принцип *усиления взаимосвязей субъективными мнениями*.

Особенность третьей ситуации в том, что исследуемые объекты имеют не только описание в виде набора показателей, но и наименования. Например, это могут быть города, районы области, фирмы, товары и т.п. Предполагается, что у экспертов сложилось полное представление об объектах, и они способны, используя принцип *субъективных предпочтений* установить значимость относительно их полезности при использовании в определенных целях.

Если объектов не более двух-трех десятков, то принцип субъективных предпочтений удобно реализовать с помощью известного метода парных сравнений. Причем, сравнивать между собой следует объекты, а не их информационные описания. Это значительно упрощает работу экспертов, но только в том случае, если у них действительно имеются интегрированные представ-

ления об объектах. В этом случае значения бинарной переменной определяются по результатам экспертного опроса, отраженным в матрице парных сравнений. Объектам, которые получили больше половины предпочтений, приписывается 1, а получившим меньше половины — 0. Одновременно с формированием массива значений бинарной переменной можно, руководствуясь первым принципом, осуществить усиление предполагаемых связей.

6. ПРОВЕРКА СОГЛАСОВАННОСТИ РЕЗУЛЬТАТОВ МОДЕЛИРОВАНИЯ ЭКСПЕРТНЫХ ПРЕДПОЧТЕНИЙ

Псевдовыборка, сформированная одним из вышеописанных способов, позволяет построить регрессионную модель, отражающую механизм субъективного выбора интересующих нас альтернатив. Этот механизм может отражать как индивидуальную точку зрения, так и групповое мнение, в зависимости от того, какая процедура использовалась при формировании выборки. Чтобы подобную модель использовать в практической деятельности, нужна уверенность в надежности получаемых с ее помощью результатов анализа и прогноза. Надежность при использовании подобного подхода рождается только из оценок, свидетельствующих о компетентности экспертов и согласованности их мнений.

Сразу согласимся с тем, что эксперт имеет право на собственное мнение, отличное от других экспертов. На наш взгляд, оригинальность в оценках не снижает уровень компетентности, несмотря на то, что есть точка зрения [1], в соответствии с которой оценка компетентности тем выше, чем ближе индивидуальные оценки к групповым. С этой точкой зрения можно согласиться только при условии, что групповые оценки совпадают или почти совпадают с истинными. Вполне возможна ситуация, когда вопреки нашим ожиданиям групповые оценки далеки от истинных значений, и заключение о компетентности на основе близости индивидуальной оценки к групповой теряет смысл.

Независимо от сделанного замечания и не взамен общепринятого подхода, введем в рассмотрение еще одну составляющую

компетентности. Оценка этой составляющей основана на положении, смысл которого в том, что эксперт не может быть компетентным, если его оценки противоречивы. Возникает естественный вопрос, как оценить уровень противоречивости экспертных суждений.

Содержательный смысл решаемой здесь задачи позволяет предложить следующий подход. Сформированная экспертами псевдовыборка представляет собой данные, которые, по сути, сгенерированы в соответствии с представлением эксперта о взаимосвязи предпочтений бинарного выбора с факторным пространством. Если модель хорошо подгоняется к этому набору данных, т.е. имеет достаточно высокий индекс отношения правдоподобия (коэффициент Макфаддена), то следует признать, что суждения эксперта не противоречивы, ему действительно удалось сформировать модель собственных суждений и поэтому он может быть отнесен к группе компетентных специалистов.

Для этих же целей можно использовать энтропийный коэффициент предсказывающих возможностей построенной модели. Этот коэффициент имеет следующий вид:

$$H = \frac{1}{n} \left[- \sum_{i=1}^n F(x_i, \mathbf{b}) \log_2 F(x_i, \mathbf{b}) - \sum_{i=1}^n (1 - F(x_i, \mathbf{b})) \log_2 (1 - F(x_i, \mathbf{b})) \right]. \quad (21)$$

Для случая, когда расчетные значения модели в точности совпадают с фактическими значениями дихотомической переменной, значение коэффициента равно нулю, т.е. во всех случаях модель обеспечивает безошибочный выбор, не оставляя место для сомнений в пользу альтернативы. Максимальное значение коэффициента ($H = 1$) получается в случае, когда для всех рассмотренных ситуаций $\hat{y}_i = F(x_i, \hat{\mathbf{b}}) = 0,5$. Во всех остальных случаях $0 < H < 1$. Величина энтропийного коэффициента показывает средний уровень неопределенности в каждом случае бинарного выбора, осуществляемого с помощью модели.

Случай, когда построенная модель оказывается неадекватной, свидетельствует о некомпетентности эксперта, так как его

взгляды на исследуемую проблему были противоречивы. Механизм, трансформирующий противоречивость взглядов в неадекватность модели, можно представить следующей схемой рассуждений. Противоречивость является следствием неуверенности эксперта в собственных суждениях. В свою очередь, неуверенность нарушает логику, которой он должен придерживаться в процессе реализации своих предпочтений на выборочном множестве. В силу этого в псевдовыборку попадают данные, слабо согласованные со значениями бинарной переменной, что исключает возможность восстановления той зависимости, которой якобы пользовался эксперт.

Таким образом, действительно модель имеет плохую точность в тех случаях, когда псевдовыборка формировалась некомпетентным образом. Поэтому уровнем адекватности модели имеет смысл измерять компетентность эксперта. Чем выше уровень адекватности (коэффициент Макфаддена) или ниже остаточная энтропия (энтропийный коэффициент), тем компетентней действовал эксперт, распределяя свои предпочтения между наблюдениями выборочной совокупности.

Переходя к рассмотрению вопроса о согласованности оценок, сразу заметим, что, как и в случае прямого экспертного опроса, в рассматриваемом подходе действует тезис — доверие к групповым оценкам выше, чем к индивидуальным. Действие тезиса возможно только при согласованности индивидуальных оценок. Следовательно, не любая групповая оценка обладает высокой надежностью. На первый взгляд может показаться, что для проверки согласованности можно воспользоваться двумя подходами: в соответствии с первым проверять согласованность до построения модели, в соответствии со вторым — после построения по результатам моделирования.

От первого подхода сразу нужно отказаться. Дело в том, что псевдовыборка, сформированная с ориентиром на согласованное мнение экспертов, не всегда гарантирует построение адекватной модели. Это бывает в том случае, когда согласованными оказались мнения некомпетентных экспертов. Подобная ситуация не возникает в задачах

прямого экспертного оценивания. В них компетентность оценивается либо экзогенно, и тогда она не связана с результатами опроса, либо в зависимости от того, насколько соответствующее индивидуальное мнение похоже на групповое.

Второй подход представляет собой авторское решение задачи, в рамках которой проверяется согласованность экспертных суждений. Как и в случае прямого экспертного оценивания в предлагаемом подходе предусматривается проверка согласованности мнений двух экспертов и проверка согласованности мнений всей группы экспертов, принявших участие в экспертизе. Обе проверки основаны на одной и той же идее. Смысл этой идеи в том, что эксперты с близкими мнениями распределяют свои предпочтения по выборочной совокупности таким образом, что полученные псевдовыборки обеспечивают построение почти идентичных моделей. Таким образом, проверка согласованности сводится к статистической проверке значимости уровня идентичности. Выполнить такую проверку можно несколькими способами.

На наш взгляд, наиболее приемлемым следует считать способ, который позволяет не только оценить статистическую значимость, но и получить содержательно интерпретируемую величину, характеризующую уровень идентичности моделей и, следовательно, уровень согласованности экспертов. В качестве такой величины удобно использовать коэффициент Юла, который измеряет тесноту связи между двумя дихотомическими переменными.

Формально с помощью этого коэффициента мы можем оценить тесноту связи между предикторными возможностями двух моделей бинарного выбора. Для этого используется таблица сопряженности 2×2 .

Правило заполнения таблицы сопряженности довольно простое. В верхней левой клеточке стоит число случаев, когда предсказания по обеим моделям совпадали и были равны 1, в нижней правой — число случаев, когда обе модели предсказали 0. В остальных клеточках стоит число несовпадающих предсказаний. Коэффициент Юла рассчитывается по формуле

$$q_{12} = \frac{ad - bc}{ad + bc}. \quad (22)$$

При полном совпадении предсказанных значений $q_{12} = 1$ и мы наблюдаем случай, когда мнения экспертов идентичны, при $q_{12} = -1$ мнения экспертов противоположны, а при $q_{12} = 0$ — независимы. Чем ближе значение коэффициента к 1, тем выше уровень согласованности экспертных мнений.

Для проверки групповой согласованности нет подходящего измерителя. Но можно предложить двухэтапную процедуру.

На первом этапе для каждой пары экспертов вычисляется коэффициент сопряженности Юла q_{ij} , и все эксперты делятся на две группы. В первую группу включаются только те эксперты, предсказанные значения по моделям которых имеют положительную связь между собой. Из коэффициентов сопряженности этой группы формируется матрица $Q = \|q_{ij}\|$.

Матрица Q обладает всеми свойствами, необходимыми для того, чтобы с помощью обычной итерационной процедуры вычислить максимальное собственное значение λ , которое является действительным числом. Тогда в качестве меры, определяющей уровень согласованности экспертов, можно использовать величину

$$L = \frac{\lambda - 1}{\text{tr } Q - 1} = \frac{\lambda - 1}{m - 1}, \quad (23)$$

в знаменателе которой стоит след матрицы без единицы.

Так определенный коэффициент согласованности будем называть характеристическим. Он равен 1, когда между всеми экспертами группы наблюдается абсолютное согласие, и равен 0, если результаты экспертного опроса статистически независимы.

С отбракованной на первом этапе группой экспертов поступают точно так же. Окончательно групповая оценка строится только для группы экспертов, имеющих согласованные мнения. Для этого все псевдовыборки объединяются в одну, по данным которой строится модель, отражающая групповое экспертное мнение. Ее и рекомендуется использовать в прогнозных расчетах.

ВЫВОДЫ

Предложенный подход позволяет строить модели, в которых через субъективные предпочтения экспертов отражены тенденции, не успевшие проявиться в статистических наблюдениях. Результаты такого моделирования удобно использовать для решения сложных прогнозных задач современ-

ного бизнеса в ситуациях, когда ощущается недостаток фактографической информации.

ЛИТЕРАТУРА

1. *Миркин, Б.Г.* Проблема группового выбора / Б. Г. Миркин. — М. : Наука, 1974. — 256 с.
2. *Green, W.H.* Econometric Analysis, 4th ed. / W. H. Green. — New York : Macmillian Publishing Company, 2000. — 1004 p.