

И.М. Никольский¹, К.К. Фурманов²

ИССЛЕДОВАНИЕ ТОЧНОСТИ РАНЖИРОВОК ПОКАЗАТЕЛЯ ЭФФЕКТИВНОСТИ В МОДЕЛИ СТОХАСТИЧЕСКОЙ ГРАНИЦЫ*

Введение

Измерение эффективности производства методами анализа данных – задача, представляющая как практический, так и академический интерес. С практической точки зрения, выявление неэффективных предприятий – одна из задач государственного регулирования. Методы анализа данных для измерения эффективности производства применяются регулирующими органами Австралии (см. [1]) и в ряде европейских стран (см. [2]).

О международном академическом интересе к данной тематике свидетельствует журнал «Journal of Productivity Analysis» с высокими наукометрическими показателями, целиком посвящённый исследованиям производственной эффективности. В отечественной академической литературе представлен ряд исследований эффективности предприятий различных отраслей [3], [4], [5] и целых регионов [6], [7] а также работы, посвящённые математико-статистическим аспектам анализа [8], [9], [10].

Спектр разработанных на данный момент подходов к анализу неэффективности достаточно широк: от метода анализа оболочки данных (Data Envelopment Analysis, DEA), впервые предложенного в [11], до регрессионных моделей границы производственных возможностей, в частности, модели стохастической границы [12].

Особенность модели стохастической границы – разделение случайной ошибки регрессии на две компоненты: неэффективность и «случайный шок». К сожалению, из-за этой особенности невозможно получить состоятельные оценки эффективности отдельных производителей (так называемых «субъектов принятия решения»), несмотря на состоятельность оценок усреднённой производственной границы по отрасли в целом.

В настоящей статье представлены результаты исследования согласованности ранжировок производителей, полученных на основании

¹ Москва, ф-т ВМК МГУ имени М.В.Ломоносова, доцент, e-mail: inikolsky@cs.msu.ru.

² Москва, ЦЭМИ РАН, лаборатория прикладной эконометрики, старший научный сотрудник, e-mail: furmach@rbcmail.ru.

модели стохастической границы. Исследования проводились с помощью серии статистических экспериментов над синтетическими случайными данными с известными ранжировками. Было показано, что даже на больших выборках точная ранжировка производителей возможна лишь в случае, когда дисперсия компоненты неэффективности превосходит дисперсию случайных шоков.

Модель стохастической границы и её оценивание

Рассмотрим выборку, представляющую собой ряд наблюдений за субъектами принятия решения (Decision Making Units, DMU). В типичных приложениях субъекты принятия решения – это фирмы-производители, эффективность которых представляет интерес для исследователя. Каждый субъект характеризуется величиной выпуска Y и затратами факторов производства $X_1, \dots, X_k$ таким образом, что наблюдаемые в данных значения выпуска определяются следующей регрессионной моделью:

$$\ln Y_i = f(X_{1,i}, \dots, X_{k,i}) + v_i - u_i,$$

где Y_i – значение выпуска Y для субъекта i ; $X_{j,i}$ – значение фактора j для субъекта i , где $j = 1, \dots, k$; v_i – случайная величина, отражающая индивидуальный сдвиг границы производственных возможностей («случайный шок»); u_i – неотрицательная случайная величина, отражающая неэффективность субъекта i . В оригинальной версии модели, предложенной [12], случайный шок нормально распределён ($v_i \sim N(0, \sigma_v^2)$), а компонента неэффективности имеет распределение полунормальное $u_i \sim N^+(0, \sigma_u^2)$ или показательное (экспоненциальное) $u_i \sim \text{Exp}(\sigma_u^{-1})$. Неизвестные параметры σ_v^2 и σ_u^2 подлежат оцениванию наряду с параметрами границы производственных возможностей. В дальнейшем получили распространение и альтернативные спецификации модели стохастической границы, в которых неэффективность имеет усечённое нормальное или гамма-распределение, при этом случайные шоки остаются нормальными.

Функция $f(X_1, \dots, X_k)$ задаёт логарифм максимального достижимого выпуска при заданных значениях факторов производства в случае однородных субъектов, при этом $\exp(f(X_1, \dots, X_k))$ – граница производственных возможностей в этом случае. Распространены два варианта функциональной формы модели. Первый из них, более простой, соответствует производственной функции Кобба–Дугласа и имеет вид

$$f(X_1, \dots, X_k) = \beta_0 + \sum_{j=1}^k \beta_j \ln X_j,$$

где $\beta_0, \dots, \beta_k$ – подлежащие оцениванию параметры модели. При этом коэффициенты $\beta_1, \dots, \beta_k$ при факторах производства равны эластичностям выпуска по соответствующим факторам. Второй вариант содержит большее количество слагаемых:

$$f(X_1, \dots, X_k) = \beta_0 + \sum_{j=1}^k \beta_j \ln X_j + \sum_{j=1}^k \gamma_j (\ln X_j)^2 + \\ + \sum_{j=1}^k \sum_{m=1}^{j-1} \delta_{j,m} (\ln X_j) (\ln X_m),$$

где $\beta_0, \dots, \beta_k, \gamma_1, \dots, \gamma_k, \delta_{1,1}, \dots, \delta_{k,k-1}$ – оцениваемые параметры. Это так называемая транслог-функция, обычно рассматриваемая как приближение для логарифма неизвестной производственной функции.

Подробный обзор практически используемых моделей стохастической границы представлен в статье [13] и книге [14]. В настоящей статье мы рассматриваем только варианты, задаваемые уравнением (1) с нормальными шоками и случайными неэффективностями u_i , имеющими полунормальное или показательное распределение.

Параметры модели стохастической границы, как правило, оцениваются методом максимального правдоподобия, что позволяет для произвольного набора значений факторов рассчитывать оценку достижимой границы $\hat{f}(X_1, \dots, X_k)$. При максимизации вместо значений σ_v^2 и σ_u^2 подбираются их логарифмы, что гарантирует положительность дисперсий. Процедуры оценивания реализованы в некоторых программах статистического анализа, таких как Stata и R (пакет frontier).

Максимум функции правдоподобия может не существовать – так будет, например, если в данных выборки распределение суммарной случайной составляющей $v_i - u_i$ окажется скошенным вправо (при выполнении предпосылок о виде распределения величина $v_i - u_i$ скошена влево, но даже в искусственно сгенерированной выборке правый хвост распределения может оказаться тяжелее левого).

При успешной максимизации после оценивания параметров модели рассчитываются суммарные логарифмические отклонения выпуска от границы $f(X_1, \dots, X_k)$:

$$\widehat{v_i - u_i} = \ln Y_i - \hat{f}(X_{1,i}, \dots, X_{k,i}).$$

Отсюда получаются оценки случайных величин u_i для каждого субъекта принятия решения:

$$\hat{u}_i = E(u_i | \{v_i - u_i = \widehat{v_i - u_i}\}).$$

В качестве меры неэффективности субъекта i может использоваться как само значение $\hat{u}_i$, так и величина $e^{-\hat{u}_i}$, лежащая в пределах от 0 (абсолютная неэффективность) до 1 (полная эффективность).

Отметим, что оценки параметров границы производственных возможностей и параметров распределений случайных компонент v_i и u_i обладают обычными свойствами оценок метода максимального правдоподобия – в частности, состоятельностью и асимптотической эффективностью. При этом оценки $\hat{u}_i$ состоятельными *не являются* – никакое увеличение выборки не позволяет устранить неопределённость, свойственную отдельному субъекту. Точное оценивание даже асимптотически возможно лишь при отсутствии случайного шока ($\sigma_v^2 = 0$), когда производственная граница едина для всех субъектов и величина u_i совпадает с наблюдаемым отклонением от этой границы. В этом случае модель будет безошибочно ранжировать субъекты принятия решения по эффективности, но это именно тот случай, когда нет смысла применять модель стохастической границы – именно предпосылка о наличии случайных шоков была новшеством Айгнера, Ловелла и Шмидта, отличавшим модель стохастической границы от предшествовавших моделей неэффективности (см., например, [15]). В настоящей статье мы приводим оценки точности ранжировок в зависимости от параметров случайных компонент модели σ_v^2 и σ_u^2 .

Используемые характеристики точности ранжировок

Для нашего исследования интерес представляет лишь способность модели давать оценки только относительной эффективности субъектов принятия решения, то есть возможность отделять более эффективные субъекты от менее эффективных или, что то же самое, способность верно ранжировать субъекты по эффективности. Вопрос точности определения ранга отдельного субъекта затронут в работе [16]. Мы рассматриваем точность ранжировки сразу всех субъектов.

При анализе эффективности для измерения согласованности ранжировок обычно используется коэффициент ранговой корреляции Спирмена. Он применяется как для определения чувствительности ранжировок к методу расчёта неэффективности (см. [13], [4]), так и для сопоставления истинных значений компоненты неэффективности u_i и их оценок в статистических экспериментах [17], что и требуется в настоящем исследовании.

Помимо коэффициента Спирмена мы рассчитываем коэффициент согласованности Харрелла (в англоязычной литературе часто называемый Harrell's C или C-index). Применительно к задаче измерения неэффективности его можно определить следующим образом. Пусть из

генеральной совокупности случайно отбираются пара субъектов. Обозначим истинные значения их компонент неэффективности u_1, u_2 , а оценки этих компонент соответственно $\hat{u}_1, \hat{u}_2$. Коэффициент согласованности Харрелла C задаётся выражением

$$C = P(\hat{u}_1 < \hat{u}_2 | u_1 < u_2).$$

В качестве выборочной оценки для C рассчитывается доля пар субъектов в выборке, для которых $\hat{u}_1 < \hat{u}_2$, среди всех пар, удовлетворяющих условию $u_1 < u_2$.

Первоначально коэффициент был предложен в работе [18] для измерения способностей медицинских тестов отражать состояние пациента. В работах [19], [20] рассматривается применение коэффициента для измерения прогнозных свойств регрессионных моделей.

Хотя коэффициент Харрелла на данный момент практически не применяется при анализе неэффективности, он имеет существенное преимущество перед коэффициентом Спирмена – интерпретируемость. Значение C лежит в пределах $[0;1]$ и отражает долю случаев, в которых модель верно выделяет относительно эффективный субъектов, среди всех пар субъектов в выборке. Например, значение $C = 0.5$ соответствует полной неспособности отличить более эффективный субъект от менее эффективного, а значение $C = 0.8$ означает, что в 80% случаев модель из пары субъектов правильно выбирает более эффективный. Здесь мы пренебрегаем вероятностью точного совпадения компонент неэффективности для разных субъектов, что вполне обоснованно, так как величины u_i предполагаются непрерывными.

План статистического эксперимента

Выбираются параметры эксперимента, в соответствии с которыми генерируются данные:

- число наблюдений n ;
- число объясняющих переменных k ;
- вектор средних значений объясняющих переменных (логарифмов факторов производства) μ ;
- ковариационная матрица объясняющих переменных Σ ;
- коэффициенты производственной границы $\beta_0, \dots, \beta_k$ (далее в статье приводятся результаты, полученные для производственной функции Кобба–Дугласа, о результатах, относящихся к транслог-функции коротко сказано в разделе 6).
- суммарная дисперсия случайных составляющих $\sigma^2 = D(v_i - u_i)$ и доля неэффективности в суммарной дисперсии $\theta = \frac{D(u_i)}{D(v_i - u_i)}$.

Значения σ^2 и θ связаны с параметрами регрессионной модели следующим образом. При полунормальном распределении неэффективности $D(u_i) = \sigma_u^2 \left(1 - \frac{2}{\pi}\right)$, так что

$$\sigma^2 = \sigma_v^2 + \sigma_u^2 \left(1 - \frac{2}{\pi}\right); \theta = \frac{\sigma_u^2 \left(1 - \frac{2}{\pi}\right)}{\sigma_v^2 + \sigma_u^2 \left(1 - \frac{2}{\pi}\right)}.$$

При показательном (экспоненциальном) распределении неэффективности $D(u_i) = \sigma_u^2$, так что

$$\sigma^2 = \sigma_v^2 + \sigma_u^2; \theta = \frac{\sigma_u^2}{\sigma_v^2 + \sigma_u^2}.$$

Приведённый список параметров эксперимента избыточен: вектор μ и коэффициент β_0 производственной функции не связаны с точностью оценивания параметров модели и были введены в план эксперимента для воспроизводимости результатов.

Основную роль играет доля неэффективности в общей дисперсии случайной составляющей θ и асимптотически этот параметр является единственным фактором, определяющим точность ранжировок, однако прочие параметры также брались в расчёт, так как они связаны с точностью оценок модели стохастической границы и предположительно могут играть роль на конечных выборках.

Представленные далее результаты получены для следующих значений параметров:

- характеристики распределения объясняющих переменных:
 $k = 2, \mu = \begin{pmatrix} 5.76 \\ 7.45 \end{pmatrix}, \Sigma = \begin{pmatrix} 0.43 & 0.54 \\ 0.54 & 0.94 \end{pmatrix};$
- характеристики распределения случайных компонент:
 $\sigma^2 = D(v_i - u_i) = 0.032$, доля неэффективности θ перебиралась от 0.1 до 0.9 с шагом 0.1.
- коэффициенты границы производственных возможностей:
 $\beta_0 = 1.17, \beta_1 = 0.60, \beta_2 = 0.38$.
- объёмы выборки $n = 25, 100, 1000$ и 10000 .

Указанные значения (кроме перебиравшихся n и θ) соответствуют хрестоматийному примеру применения модели стохастической границы из книги [21], опирающемуся на данные исследования предприятий металлургической отрасли США 1957 года. Эти же данные, но с дополнениями, использовались в основополагающей статье по анализу стохастической границы [12]. Результаты, полученные для других значений параметров эксперимента, описаны вкратце в разделе 6.

Схема генерации данных в экспериментах следующая:

1. Генерируются n случайных значений вектора объясняющих переменных $(X_{1,i}, \dots, X_{k,i}), i = 1, \dots, n$ в соответствии с многомерным нормальным распределением с вектором математических ожиданий μ и ковариационной матрицей Σ . Предполагается, что переменные $(X_{1,i}, \dots, X_{k,i})$ отражают уже логарифмированные значения затрат факторов производства.
2. Генерируются значения шоков $v_1, \dots, v_n$, так что $v_i \sim N(0, \sigma_v^2)$, где $\sigma_v^2 = (1 - \theta)\sigma^2$.
3. Генерируются значения компонент неэффективности $u_1, \dots, u_n$, $u_i \sim N^+(0, \sigma_u^2)$, где $\sigma_u^2 = \frac{\theta\sigma^2}{1 - \frac{\pi}{2}}$ для полунормального распределения и $\sigma_u^2 = \theta\sigma^2$ для экспоненциального.
4. Генерируются значения объясняемой переменной – логарифма выпуска: $\ln Y_i = \beta_0 + \sum_{j=1}^k \beta_j X_j + v_i - u_i$.

При генерации предполагалось, что наблюдения за разными субъектами независимы.

По каждой сгенерированной выборке методом максимального правдоподобия оцениваются параметры модели стохастической границы и компоненты неэффективности $\hat{u}_1, \dots, \hat{u}_n$. Если при максимизации правдоподобия сходимость не достигается за 200 итераций метода Ньютона–Рафсона, фиксируется отсутствие сходимости³. Иначе рассчитываются коэффициент корреляции Спирмена и коэффициент согласованности Харрелла между истинными (сгенерированными) значениями u и их оценками $\hat{u}$.

Для каждой пары перебираемых значений n и θ эксперимент проводился 1000 раз, в результате каждого опыта программа возвращала либо отметку о недостижении сходимости, либо значения коэффициентов Спирмена и Харрелла. По каждой серии из 1000 повторений рассчитывались средние значения коэффициентов корреляции и квантили порядков 0.05 и 0.95 для получения представления о разбросе значений коэффициентов.

Все расчёты проводились в программе Stata версии 11.0 [22].

Результаты эксперимента

В данном разделе представлены результаты статистических экспериментов на случайных данных с полунормальным (табл.1) и экспоненциальным (табл. 2) распределением. Наряду с долей опытов,

3 Первоначально был поставлен предел в 16000 итераций, но обнаружилось, что если сходимость не достигается на 200 итерациях, она не достигается и после.

закончившихся сходимостью процедуры максимизации правдоподобия, и средних значений коэффициентов Спирмена и Харрелла представлены квантили коэффициента Харрелла порядков 0.05 и 0.95, дающие представление о разбросе этого коэффициента. Результаты для выборок объема 1000 опущены, так как они практически не отличаются от случая 10000 наблюдений за исключением разницы в разбросе коэффициента согласованности.

В таблицах используются следующие обозначения: n – объем выборки; θ – доля дисперсии неэффективности в суммарной дисперсии случайной составляющей; $\bar{\rho}$ – среднее значение коэффициента ранговой корреляции Спирмена между истинными и оценёнными значениями компоненты неэффективности; $\bar{C}$ – среднее значение коэффициента согласованности Харрелла между истинными и оценёнными значениями компоненты неэффективности; $C_{0.05}$ и $C_{0.95}$ – квантили порядков 0.05 и 0.95 для коэффициента Харрелла. Под «сходимостью» понимается доля опытов, когда был достигнут максимум функции правдоподобия.

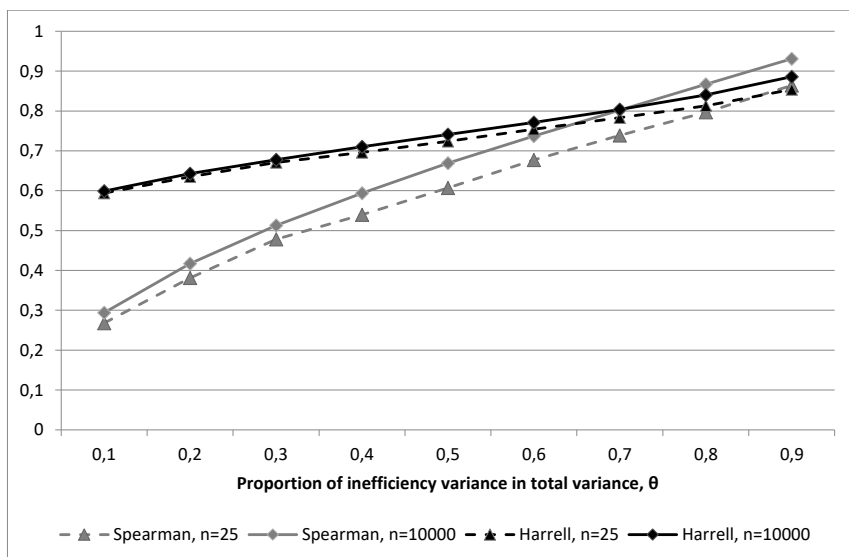


Рис. 1. Коэффициенты корреляции Спирмена и Харрелла между истинными показателями неэффективности и их оценками в зависимости от доли дисперсии компоненты неэффективности в суммарной дисперсии случайной составляющей, выборки объема 25 и 10000 наблюдений.

Рассмотрим подробнее результаты, полученные для полунормального распределения неэффективности. Как можно видеть из табл.1, согласованность модельных ранжировок с истинными увеличивается с ростом доли дисперсии неэффективности в суммарной дисперсии случайной составляющей, при этом объём выборки оказывает довольно скромное влияние на коэффициенты согласованности.

Из рис. 1 видно, что модель способна точно определять относительно эффективные субъекты лишь в случае, когда разброс случайной составляющей достигается в основном за счёт различий субъектов в эффективности, а вклад случайных шоков мал. Так, если 90% дисперсии случайной составляющей приходится на долю неэффективности, то коэффициент Спирмена на больших выборках равен 0.87–0.88, а коэффициент Харрелла – 0.85–0.86, то есть доля ошибок при попарных сравнениях субъектов по эффективности составляет примерно 15%.

Если же основной компонентой случайной составляющей является шок, то достижение высокой согласованности истинных и модельных ранжировок достижимо только благодаря случайности на малых выборках, когда коэффициенты согласованности имеют большую дисперсию, при этом возможен и обратный результат: модель даёт ранжировки скорее противоположные, чем согласующиеся с действительными. Так, при $\theta = 0.1$ и $n = 25$ коэффициент Харрелла оказывался ниже 0.5 более чем в 5% случаев ($C_{0.05} = 0.487$, см. табл. 1) – в этих случаях попарные сравнения модельных оценок неэффективности чаще приводили к ошибочным выводам, чем к правильным.

Серия экспериментов, проведённых для модели с экспоненциальным распределением компонент неэффективности $u_1, \dots, u_n$, показала значения коэффициентов Харрелла и Спирмена, меньшие примерно на 0.01–0.03 по сравнению с соответствующими значениями для модели с полунормальным распределением, причём расхождения не уменьшались при увеличении выборки. Наибольшие различия выявлялись в случаях, когда доли дисперсии случайных шоков и неэффективности были близки друг к другу. В целом, результаты, полученные для полунормального и для экспоненциального распределения неэффективности, схожи.

Проведённые эксперименты показывают, что модель стохастической границы способна приводить к точным ранжировкам только в случае, когда дисперсия неэффективности больше дисперсии случайных шоков ($\theta > 0.5$).

$n = 25$					
θ	0.1	0.2	0.3	0.4	0.5
Сходимость	0.965	0.963	0.963	0.941	0.936
$\bar{\rho}$	0.268	0.381	0.478	0.540	0.607
$\bar{C}$	0.594	0.635	0.671	0.696	0.724
$C_{0.05}$	0.487	0.523	0.564	0.593	0.623
$C_{0.95}$	0.707	0.737	0.767	0.790	0.810
θ	0.6	0.7	0.8	0.9	
Сходимость	0.937	0.907	0.908	0.844	
$\bar{\rho}$	0.677	0.739	0.797	0.864	
$\bar{C}$	0.754	0.783	0.813	0.854	
$C_{0.05}$	0.666	0.690	0.732	0.774	
$C_{0.95}$	0.833	0.857	0.880	0.910	
$n = 100$					
θ	0.1	0.2	0.3	0.4	0.5
Сходимость	1.000	1.000	1.000	1.000	1.000
$\bar{\rho}$	0.289	0.408	0.500	0.585	0.655
$\bar{C}$	0.598	0.641	0.675	0.708	0.737
$C_{0.05}$	0.546	0.591	0.628	0.662	0.696
$C_{0.95}$	0.650	0.688	0.721	0.751	0.777
θ	0.6	0.7	0.8	0.9	
Сходимость	1.000	0.998	0.997	0.988	
$\bar{\rho}$	0.727	0.789	0.851	0.918	
$\bar{C}$	0.769	0.800	0.834	0.880	
$C_{0.05}$	0.730	0.763	0.801	0.853	
$C_{0.95}$	0.807	0.831	0.864	0.902	
$n = 10000$					
θ	0.1	0.2	0.3	0.4	0.5
Сходимость	1.000	1.000	1.000	1.000	1.000
$\bar{\rho}$	0.294	0.417	0.513	0.594	0.669
$\bar{C}$	0.599	0.643	0.678	0.710	0.741
$C_{0.05}$	0.590	0.635	0.674	0.705	0.736
$C_{0.95}$	0.605	0.648	0.683	0.714	0.745
θ	0.6	0.7	0.8	0.9	
Сходимость	1.000	1.000	1.000	1.000	
$\bar{\rho}$	0.737	0.802	0.867	0.931	
$\bar{C}$	0.771	0.804	0.840	0.886	
$C_{0.05}$	0.768	0.800	0.837	0.884	
$C_{0.95}$	0.775	0.807	0.843	0.888	

Табл.1 Результаты экспериментов с полунормальным распределением компоненты неэффективности.

$n = 25$					
θ	0.1	0.2	0.3	0.4	0.5
Сходимость	0.604	0.652	0.714	0.771	0.821
$\bar{\rho}$	0.265	0.348	0.430	0.508	0.555
$\bar{C}$	0.593	0.624	0.654	0.685	0.703
$C_{0.05}$	0.480	0.510	0.540	0.577	0.597
$C_{0.95}$	0.700	0.727	0.753	0.778	0.800
θ	0.6	0.7	0.8	0.9	
Сходимость	0.876	0.916	0.956	0.982	
$\bar{\rho}$	0.609	0.678	0.746	0.829	
$\bar{C}$	0.727	0.757	0.789	0.835	
$C_{0.05}$	0.617	0.657	0.696	0.757	
$C_{0.95}$	0.820	0.840	0.873	0.897	
$n = 100$					
θ	0.1	0.2	0.3	0.4	0.5
Сходимость	0.662	0.749	0.874	0.945	0.977
$\bar{\rho}$	0.267	0.366	0.454	0.521	0.589
$\bar{C}$	0.591	0.626	0.658	0.684	0.711
$C_{0.05}$	0.538	0.568	0.606	0.636	0.661
$C_{0.95}$	0.645	0.679	0.707	0.733	0.757
θ	0.6	0.7	0.8	0.9	
Сходимость	0.998	0.998	0.999	1.000	
$\bar{\rho}$	0.653	0.723	0.793	0.874	
$\bar{C}$	0.739	0.770	0.805	0.852	
$C_{0.05}$	0.692	0.727	0.763	0.818	
$C_{0.95}$	0.782	0.808	0.841	0.881	
$n = 10000$					
θ	0.1	0.2	0.3	0.4	0.5
Сходимость	0.994	1.000	1.000	1.000	1.000
$\bar{\rho}$	0.266	0.373	0.458	0.532	0.600
$\bar{C}$	0.590	0.628	0.659	0.687	0.714
$C_{0.05}$	0.584	0.622	0.654	0.682	0.709
$C_{0.95}$	0.595	0.633	0.664	0.692	0.718
θ	0.6	0.7	0.8	0.9	
Сходимость	1.000	1.000	1.000	1.000	
$\bar{\rho}$	0.665	0.732	0.804	0.885	
$\bar{C}$	0.742	0.772	0.808	0.856	
$C_{0.05}$	0.737	0.768	0.804	0.853	
$C_{0.95}$	0.746	0.776	0.811	0.859	

Табл.2 Результаты экспериментов с показательным (экспоненциальным) распределением компоненты неэффективности.

Чувствительность результатов к выборам параметрам эксперимента

Для исследования чувствительности данные генерировались для наборов параметров, отличных от приведённых в разделе 4. В частности:

- дисперсия объясняющих переменных уменьшалась и увеличивалась до 100 раз;
- корреляция между объясняющими переменными изменялась от 0 до 0,999;
- суммарная дисперсия случайных компонент уменьшалась и увеличивалась до 100 раз;
- коэффициенты границы производственных возможностей $\beta_0, \beta_1, \beta_2$ уменьшались и увеличивались до 10 раз;
- число объясняющих переменных варьировалось от 1 до 5 (что позволяет распространить полученные результаты и на случай границы производственных возможностей, задаваемой транслог-функцией).

Единственным параметром, который заметно и систематически влиял на результаты, оказалось число объясняющих переменных. Увеличение числа переменных снижает согласованность ранжировок на малых выборках, причём снижение тем больше, чем больше доля неэффективности в общей дисперсии. Так, при $k = 5$ и $n = 25$ средний коэффициент Харрелла равнялся 0.567 для $\theta = 0.1$ и 0.812 для $\theta = 0.9$ (соответствующие значения при $k = 2$ и том же объёме выборки – 0.594 и 0.854 соответственно). При увеличении объёма выборки параметр k перестаёт играть роль. В случае $n = 100$ наибольшее расхождение в значении коэффициента Харрелла между моделями с двумя и с пятью объясняющими переменными составило 0.013. Для выборок большего объёма сколь-либо существенного влияния параметра k на результаты не обнаружено.

При изменении прочих параметров заметно менялась доля случаев недостижения сходимости, но при сходимости значения коэффициента Харрелла отличались от представленных в табл.1 менее чем на 0.02 на малых выборках ($n = 25$), и мы не имеем оснований считать, что эти изменения не случайны. При увеличении выборок различия становились ещё меньше.

Заключение

Низкая чувствительность коэффициентов согласованности ранжировок между истинными и модельными показателями неэффективности к параметрам эксперимента свидетельствует, что даже на небольших выборках точность оценок эффективности определяется в

основном соотношением дисперсий неэффективности и случайных шоков. Приведённые в табл. 1 и 2 значения коэффициентов Спирмена и Харрелла могут быть использованы исследователями как ориентировочные характеристики точности модели, применённой к разным выборкам, так как оценки дисперсий могут быть без труда получены для произвольного набора данных (если только достигается максимум функции правдоподобия), а остальные характеристики данных играют относительно малую роль.

Результаты проведённого эксперимента свидетельствуют, что точные оценки эффективности отдельных субъектов возможны только в случае, когда в суммарной дисперсии случайной составляющей в модели стохастической границы основная доля приходится на неэффективность субъектов. Эти случаи могут быть легко выявлены, так как оценивание модели предполагает оценивание дисперсий компонент случайной составляющей.

Литература

1. Australian Competition and Consumer Commission. (2012). Benchmarking Opeх and Capex in Energy Networks. ACCC/AER Working Paper Series, No. 6, May 2012.
2. Schweinsberg A., Stronzik M., Wissner M. (2011). Cost Benchmarking in Energy Regulation in European Countries. WIK-Consult (Final Report)
3. Щетинин Е.И., Назруллаева Е.Ю. (2012). Производственный процесс в пищевой промышленности: взаимосвязь инвестиций в основной капитал и технической эффективности. Прикладная эконометрика, 28(4): 63–84.
4. Илатова И.Б. (2015). Динамика совокупной факторной производительности и её компонентов на примере российской отрасли, производящей пластмассовые изделия. Прикладная эконометрика, 38(2): 21–40.
5. Bessonova E., Tsvetkova A. (2019). Productivity convergence trends within Russian industries: Firm-level evidence. Bank of Russia Working Paper Series wps 51. Bank of Russia.
6. Макаров В., Айвазян С., Афанасьев М., Бахтизин А., Нанавян А. (2016). Моделирование развития экономики региона и эффективность пространства инноваций. Форсайт, 10(3): 76–90
7. Айвазян С.А., Афанасьев М.Ю., Кудров А.В. (2016). Модели производственного потенциала и оценки технологической эффективности регионов РФ с учётом структуры производства. Экономика и математические методы, 52(1): 28–44

8. Айвазян С.А., Афанасьев М.Ю., Руденко В.А. (2014). Исследование зависимости случайных составляющих стохастической производственной функции при оценке технической эффективности. Прикладная эконометрика, 34(2): 3–18.
9. Rudenko V.A., Aivazyan S.A., Afanasyev M.Y. (2017). Specification of a Stochastic Production Function Model in the Extended Class of Stochastic Frontier Models. Modeling of Artificial Intelligence, 4(1): 21–28.
10. Rudenko V.A. (2018). Specification Scheme of the Stochastic Production Function for Assessment of Technical Efficiency of the Regions in the Russian Federation. Russian Journal of Mathematical Research, Series A, 4: 38–47.
11. Farrell M. (1957). The Measurement of Productive Efficiency. Journal of the Royal Statistical Society, Series A, General, 120: 253–281.
12. Aigner D., Lovell C.A.K., Schmidt P. (1977). Formulation and estimation of stochastic frontier function models. Journal of Econometrics, 6: 21–37.
13. Малахов Д.И., Пильник Н.П. (2013). Методы оценки показателя эффективности в моделях стохастической производственной границы. Экономический журнал Высшей Школы Экономики, 17(4): 660–686.
14. Kumbhakar S.C., Lovell C.A.K. (2003). Stochastic Frontier Analysis. Cambridge University Press.
15. Winsten C. (1957). Discussion on Mr. Farrell's Paper. Journal of the Royal Statistical Society, Series A, General, 120: 282–284.
16. Horrace W.C., Seth R.-S., Wright I. (2015). Expected efficiency ranks from parametric stochastic frontier models. Empirical Economics, 48(2): 829–848.
17. Feng D., Wang C., Zhang X. (2019). Estimation of inefficiency in stochastic frontier models: a Bayesian kernel approach. Journal of Productivity Analysis, 51(1): 1–19.
18. Harrell F.E. Jr., Califf R.M., Pryor D.B., Lee K.L., Rosati R.A. (1982). Evaluating the yield of medical tests. Journal of the American Medical Association, 247(18): 2543–2546.
19. Newson R.B. (2010) Comparing the predictive powers of survival models using Harrell's C or Somers' D. The Stata Journal, 10(3): 339–358.
20. Rumyantseva E.V., Furmanov K.K. (2021). Using out-of-sample Cox–Snell residuals in time-to-event forecasting. Business Informatics, 15(1): 7–18.
21. Greene W.H. (2018). Econometric Analysis. 8th edition. Pearson.
22. www.stata.com.