

16

FLUID VULNERABILITY THEORY: A COGNITIVE APPROACH TO UNDERSTANDING THE PROCESS OF ACUTE AND CHRONIC SUICIDE RISK

M. DAVID RUDD

It is now well accepted that the assessment, management, and treatment of suicidality in clinical practice is one of the most challenging and stressful tasks for any clinician (Jobes, 1995). The literature in suicidology routinely differentiates among treatment, treatment outcome, and risk assessment (e.g., Rudd, Joiner, & Rajab, 2000), with no clear theoretical link across the three areas. Additionally, there has been limited work addressing content versus process issues in each area specific to suicide risk assessment. The current theory being offered focuses specifically on the risk assessment process, not treatment outcome. Furthermore, its focus is not on the specific content of risk assessment (i.e., what questions to ask across what content domains). A considerable amount is known about the *content* of risk assessment (e.g., Rudd et al., 2000). This is a fairly significant departure from the routine in suicidology, but it is one I believe to be important for a number of reasons that are emphasized in this chapter.

Empirically grounded risk assessment models have emerged over the past several years, proving to be both effective and efficient in clinical practice (e.g., Joiner, Walker, Rudd, & Jobes, 1999; Rudd et al., 2000). Most of these models have emphasized the need to evaluate and respond to identifiable risk and protective factors (e.g., Clark & Fawcett, 1992; Hirschfeld & Russell, 1997; Joiner et al., 1999). Among the more commonly targeted risk factor domains are previous suicidal behavior, current suicidal symptoms, precipitant stressors, associated clinical symptoms, impulsivity, and self-control. Protective factors have routinely included social support among family and friends and previous or current treatment involvement, along with the efficacy of previous treatment efforts. Although these assessment models have led to useful clinical heuristics, they have not been driven by coherent theory—theory characterized by clearly defined constructs that lead to specific, testable hypotheses and ultimately help advance the science of clinical suicidology. Rather, the traditional approach to risk assessment has been to categorize or otherwise organize empirically supported risk factors and apply them to clinical practice, regardless of whether there was any underlying theory that would tie findings together in a meaningful way. In some ways, this approach neglects critical process issues to risk assessment, the most important of which is an understanding of the relationship between acute suicidal states and those that endure or recur over longer periods of time. This is one area of significant need in clinical suicidology, that is, theory that is specific to the variable nature of acute suicide risk and relates it to potential enduring risk over time.

There is a host of questions that beg attention. How can the emergence, persistence, resolution, and, ultimately, reemergence of suicide risk over time for a given individual be explained? How can the clinician understand the relationship between those who only think about suicide, those who make a single attempt, and those who make multiple suicide attempts over a lifetime? What is distinctive about a single specific episode of suicidality? How does one episode relate to subsequent episodes? Does one episode set the stage for subsequent episodes? If so, what is the mechanism of action? Another way of framing these questions is to address distinctions among those with suicidal ideation, those who make a single suicide attempt, and those who make multiple attempts (e.g., Rudd, Joiner, & Rajab, 1995). In other words, how is someone who makes a single attempt different from someone who makes multiple attempts? Why do some people with suicidal ideation and some who make single attempts go on to make multiple attempts, and why do some not? An emerging literature addresses what appears to be the distinctive nature and characteristics of those with multiple attempts, and this needs to be considered and integrated into the broader risk assessment literature (e.g., Rudd et al., 1995).

In some respects, this is a new kind of theory, one specific to the process of acute and chronic suicide risk, a theory grounded in identified fundamen-

tal assumptions and tied to broader models for understanding and treating suicidality. There is little debate that a considerable number of theoretical approaches exist for understanding suicidality, including philosophical, psychiatric, psychodynamic, sociological, sociocultural, and psychological (e.g., Rudd, 2004). What seems to be missing, however, is theory that is specific to the process of risk, theory that attempts to understand and explain the parameters of suicidal crises, both the intensity and duration of a single suicidal episode and how it relates to subsequent or multiple episodes. What is needed is theory that helps with questions such as the following: How does an episode of risk emerge? How long does it last? How intense or severe is it? What's the nature or course of recovery? Is the episode related to subsequent episodes? Is there a difference in the *process* of risk for different patients? Do those who attempt suicide multiple times become suicidal under different conditions and recover more slowly than those with only a single attempt? If so, what are the characteristic features of this process? Is it possible to identify process types, such as short- versus long-interval subsequent attempts? Needless to say, answers to these questions would prove invaluable in day-to-day clinical decision making. These are the kinds of questions clinicians ask with increasing frequency. How long will the suicidal state last? When acute risk resolves, should I worry or does it generate vulnerability for subsequent episodes? What steps can be taken to reduce risk effectively over both the short and long term?

Recent discussions about possible warning signs for suicide helped underscore the problem in the clinical practice of risk assessment (e.g., Rudd, 2003). Ultimately, warning signs are very different from risk factors, both conceptually and in terms of immediate and longer term clinical decision making. There is a **temporal aspect to warning signs that convey the notion of imminent risk**—that something will happen in the next few minutes or hours. For the most part, risk factors have no identified time constraint. A quick review of the literature reveals that research tying risk factors to suicidality often do so over long periods of time, with the minimum being at least 12 months (e.g., Joiner et al., 1999) and the maximum several decades. How, then, do extant findings relate to episodes of imminent risk? This is a question of considerable importance but one with no empirical data to respond. **What is needed is research that looks at risk over very brief periods of time that are of particular salience for the clinical encounter**, periods of an hour, a day, or a week. What is also needed is theory to drive this new line of research, focusing on the **process of risk over the short and longer term**. It is probably fair to say that clinicians have a good and ever-improving idea of what questions to ask in the immediate clinical setting and what clinical markers to look for and address. Clinical practice steadily improves, as science helps identify which factors are correlated to risk. In contrast, it is difficult to say that much is known about the process of risk, that is, how long an acute episode might last for a given patient, what specific factors are associ-

ated with longer duration suicidal crises or those of greater intensity, or what patients' risk status might be when the acute episode resolves. Resolution of an acute episode certainly does not resolve long-term or chronic risk. Understanding the nature of enduring or chronic risk is of critical importance in clinical practice. Being able to recognize and respond in effective fashion to chronic risk is vital to safe and effective practice.

FLUID VULNERABILITY THEORY: FUNDAMENTAL ASSUMPTIONS

Fluid vulnerability theory (FVT) is a way to understand the process of suicide risk, over both the short and longer term. It is a theory embedded in cognitive theory and therapy and the idea of the suicidal mode. I have discussed the notion of the suicidal mode at length elsewhere and do not repeat all the details here, but it is important to understand that the two ideas are very much complementary and interwoven (Rudd et al., 2000). Here I provide the necessary overview of the suicidal mode to ground FVT effectively. FVT is a way of understanding the onset of episodes of risk, that is, acute activation of the suicidal mode. Episodes of risk can be conceptualized as suicidal states, with some identifiable aspects that remit and others that endure over time, sometimes for long periods. It helps answer questions such as why people become suicidal, how long they will stay suicidal, how severe an episode will be, and whether there is high probability for another episode. At the most fundamental level, FVT is guided by the assumption that suicidal episodes are time limited. That is, they do not last an indeterminant period but endure for as long as the suicidal mode is active. In other words, the characteristic features of an individual's suicidal mode (i.e., suicidal belief system, physiological-affective symptoms, and associated behaviors and motivations) provide information on which to formulate hypotheses about the following: susceptibility to an episode of suicidality, likely triggers (i.e., precipitants), duration of an episode, and potential for future episodes (i.e., risk for chronic suicidality). In short, FVT hypothesizes that the state of suicidality, the factors that triggered the episode, and those that contribute to its severity and duration are fluid in nature and duration. That is, an individual's vulnerability to suicide is variable but nonetheless identifiable and quantifiable. A patient can be suicidal on Tuesday (i.e., an activated suicidal mode) and at low risk on Wednesday if the mode is effectively deactivated.

A few quick points about the suicidal mode need to be made here to provide the necessary foundation. First, the suicidal mode has four components or domains: the suicidal belief (cognitive) system, the affective system, the physiological system, and the behavioral (motivational) system. The four systems work in synchrony when triggered by either an internal (e.g., thought, feeling, image) or external precipitant (e.g., the loss of a relationship). The

possibility of internal triggering provides a means to understand biological vulnerability, but more about that later. The **end result is a suicidal episode or state that is characterized by specific or core cognitive themes** (i.e., unlovability, helplessness, poor distress tolerance, and perceived burdensomeness), **acute dysphoria and related physiological arousal** (i.e., Axis I symptomatology), **and associated death-related behaviors**. As discussed in more detail later, these themes can each be tied to specific core beliefs, underlying assumptions, and related automatic thoughts (i.e., consistent with Beck's three levels of cognition; Alford & Beck, 1997). First, however, I briefly examine the suicidal mode (the reader is referred to Rudd et al. [2000] for a full and detailed description).

Beck (1996) recently offered a refinement of his original cognitive therapy model in response to a growing body of empirical studies and theoretical discourse that highlighted a number of shortcomings in efforts to explain more complex theoretical constructs and related interactions and then experimentally validate them (e.g., Haaga, Dyck, & Ernst, 1991). The model is consistent with the axioms noted earlier and builds on the concept of schemas and simple linear schema processing in a number of important ways. The theory is built around the concept of the *mode*, the structural or organizational unit that contains schemas. Beck (1996) defined modes as "specific suborganizations within the personality organization [that] incorporate the relevant components of the basic systems of personality: cognitive (or information processing), affective, behavioral, and motivational" (p. 4). He went on to note that, consistent with the original theory, each system is composed of structures identified as *schemas* (e.g., affective schemas, cognitive schemas, behavioral schemas, and motivational schemas). Beck integrated the physiological system as separate but noted its unique and significant contribution to the overall functioning of the mode. Of particular importance to the concept of the mode is the previously noted issue of reciprocal determinism and synchrony of action. Beck (1996) described the mode as an "integrated cognitive–affective–behavioral network [that] produces a synchronous response to external demands and provides a mechanism for implementing internal dictates and goals" (p. 4).

The **cognitive system** is described as involving all aspects of information processing including selection of data, attentional process (i.e., meaning assignment and meaning making), memory, and subsequent recall. Incorporated within this system is the notion of the *cognitive triad*, integrating beliefs regarding self, others, and the future. **For our discussion of the suicidal mode, the representative cognitive triad, along with the associated conditional assumptions, rules, and compensatory strategies, is referred to as the suicidal belief system (SBS)**. Consistent with Alford and Beck's (1997) Axiom 6, three levels of cognition are assumed, with the majority of therapeutic efforts targeting the more conscious levels. This does not negate, however, the importance of preconscious and metacognitive processing. Additionally, it pro-

vides a means of integrating research and theory on implicit learning and tacit knowledge (e.g., Dowd & Courchaine, 1996).

The *affective system produces emotional and affective experience*. Beck (1996) noted the importance of the affective system, emphasizing its role in reinforcing adaptive behavior, through the experience of both positive and negative affect (Beck, Emery, & Greenberg, 1985). This makes both conceptual and logical sense. He went on to state that negative affective experiences serve to *focus the attention* of individuals on circumstances or situational contexts that are not in our best interest or serve to *diminish [us] in some way* (1996, p. 5). As a result, a negative valence is created for that event, situation, or experience, increasing sensitivity of the mode to being triggered or activated in the future under comparable circumstances. This helps explain low threshold for activation for some suicidal patients, as well as generalization across similar but not entirely identical situations or circumstances. For multiple attempters, then, they would not only have a lower *activation threshold* but also a broader range of internal and external triggers.

Finally, the *motivational and behavioral systems* allow for autonomic activation or deactivation of the individual for response. Although Beck noted that the motivational and behavioral systems are, for the most part, automatic in activation, they can be consciously controlled under some conditions. The *physiological system* comprises the physiological symptomatology accompanying the mode. For a *threat mode*, for example, this would include autonomic arousal, along with motor and sensory system activation. This would serve to orient the individual for action such as *fight or flight*. The synchronous and simultaneous interaction of multiple systems, and potential cognitive misinterpretation during a threat mode, leads to escalation and expansion of physical symptoms (e.g., perception of threat from panic symptoms such as “I’m having a heart attack”). Again, each system comprises structures or *schemas* specific to that system. Accordingly, the suicidal belief system comprises beliefs or schemas within each of the identified systems (i.e., affective schemas, behavioral schemas, and motivational schemas).

In short, FVT provides a means to understand the onset of an episode of acute suicidality, how one episode relates to another, how long the episode might last, and how severe it might be across each of the four mode domains. FVT essentially explains the *process* of suicidality, that is, why some people make a suicide attempt and never experience another episode and why another might make multiple attempts that span 20 to 30 years, some with brief intervals and others with very long intervals between attempts. Consistent with Litman’s (1991) notion of the *suicide zone*, *FVT predicts that suicide risk is time limited, that imminent risk cannot endure beyond periods of heightened arousal (i.e., activation across all four mode domains)*. This is not to say that chronic suicidality does not exist; it certainly does, but it is best understood as recurrent (and discrete) periods of imminent risk rather than enduring risk over time. In addition, *susceptibility to subsequent at-*

tempts can be understood as a function of vulnerability to activation of the suicidal mode or *triggering*, susceptibility that extends across all domains:

- cognitive susceptibility (to include impaired problem solving, a lack of cognitive flexibility or cognitive rigidity, cognitive distortions inherent to the suicidal belief system, etc.);
- biological susceptibility (i.e., physiological and affective symptoms); and
- behavioral susceptibility (e.g., deficient skills that cut across a broad range of areas, such as interpersonal, self-soothing, and general emotion regulation).

As mentioned earlier, FVT is guided by a number of fundamental assumptions. The **foundational assumption** is that suicidal episodes are time limited. A **second assumption** is that baseline risk varies from individual to individual. Everyone has a baseline risk level, a threshold value at which the suicidal mode is activated, a value that is set in accordance with susceptibility across each domain just mentioned. It is important to understand that this threshold value varies for each individual. For some it may be so high that the suicidal mode is unlikely ever to be activated. In other words, it may be that for some people, there are no conditions under which they would ever consider suicide. Take for example the fact that under some of the most extreme and harsh conditions (e.g., prisoners of war), certain individuals would never consider suicide an option. In this case, the suicidal belief system is characterized by thoughts such as *I would never kill myself under any conditions* or core beliefs such as *life is too precious ever to consider suicide*. For others, however, the threshold value for activation is very low, with minimal stressors, physiological–affective symptoms, or limited skills, triggering a cascade of **suicidal thoughts that tend to cluster around four primary themes**: *I'm worthless and don't deserve to live* (core belief of unlovability), *I can't fix this problem and should just die* (core belief of helplessness), *I'd rather die than feel this way* (core belief of poor distress tolerance, i.e., negative emotion is intolerable), or *everyone would be better off if I were dead* (core belief of perceived burdensomeness). As suggested here, there are **four primary core belief themes that are revealed in expressed automatic thoughts: unlovability, helplessness, poor distress tolerance, and perceived burdensomeness**. I am suggesting here that these core beliefs are potentially quite different from those routinely associated with depression and are specific to suicidality (i.e., cognitive content specificity).

Regardless of relative value across each domain of the suicidal mode, FVT predicts that everyone has a baseline risk level that is determined by historical and developmental factors. That baseline risk level predicts ease of activation or why someone might become suicidal under a given set of conditions. In short, vulnerability can be understood as a function across each domain: cognitive, affective, physiological, and behavioral. As a patient im-

proves in one area, vulnerability to subsequent episodes would be reduced. However, it is critical to keep in mind that an extremely low threshold for activation in any given area would essentially undermine or limit progress in another. For example, marked improvement in depressive symptoms and related physiological arousal could be undercut by easily activated depressive or suicidal schemas. This is generally consistent with the old adage that depressed patients can be at heightened suicide risk after neurovegetative symptoms have remitted.

From a cognitive perspective, the suicidal belief system helps determine this baseline risk. In other words, those with high versus low baseline risk levels would be characterized by very different cognitive content (i.e., again consistent with the notion of cognitive content specificity). It is likely that those experiencing chronic suicidality and making multiple attempts would endorse automatic thoughts representative of core beliefs across more themes than those who make single attempts, have suicidal ideation, or are depressed. From a cognitive theory standpoint, this is one way of understanding the importance of historical factors in most risk assessment models. For example, most risk models have incorporated the importance of historical factors such as previous attempts (e.g., Clark & Fawcett, 1992; Rudd et al., 1995), and previous psychiatric diagnoses (e.g., Tanney, 1992). The argument here is that those historical factors are characterized by cognitive content and variability exists across those that are highly vulnerable (i.e., those with high baseline risk) and those that are resilient (i.e., low baseline risk). This is not to minimize the importance of biological vulnerability but rather to emphasize that previous or multiple episodes of a mood disorder have associated cognitive content. That is, there is impact on the suicidal belief system, incorporating beliefs relevant to both self and others.

The **third assumption in FVT** states that after resolution of an acute episode an individual returns to his or her baseline risk level. This is important because baseline risk varies so much from individual to individual, as suggested earlier. If the clinician is working with a patient who has high baseline risk (i.e., the suicidal mode and acute episodes are easily triggered), resolution of an acute episode may not necessarily mean that risk has resolved to any significant degree. Accordingly, the **fourth assumption in FVT** states that those who make multiple attempts (i.e., two or more genuine suicide attempts) have higher baseline risk levels. In other words, of all identifiable groups, those individuals with multiple attempts are at greatest chronic risk. In short, repeated suicidal states have resulted in a suicidal mode that is easily triggered, with activation occurring across any of the four domains (i.e., internal and external triggering).

Several investigators have found that those with multiple attempts are at significantly higher enduring risk relative to those who make single attempts and those who have suicidal ideations (Clark & Fawcett, 1992; Rudd et al., 1996). In accordance with this assumption, FVT predicts that those

who make multiple attempts would be characterized by suicidal belief systems that have the greatest breadth and depth of beliefs across the four themes identified earlier. It is also hypothesized that these individuals would be characterized by the most severe symptoms (i.e., physiological and affective systems) and related behavioral deficits (e.g., emotion regulation and interpersonal skills). As noted earlier, emerging data support these hypotheses (Clark & Fawcett, 1992; Rudd et al., 1995).

The fifth assumption in FVT states that suicide risk is elevated by aggravating factors, which essentially are precipitant stressors (either internal or external) that cut across the four domains of the suicidal mode. This period of aggravation is time limited in nature. It is important to remember, consistent with what was noted previously, that precipitant stressors can be *internal or external* and cut across the four domains of the suicidal mode. More specifically, a patient with a recurrent major depressive disorder can have the suicidal mode triggered by a reemergence of depressive symptoms. A patient with poor interpersonal skills can have an episode triggered by an argument with a close friend. What is important is the notion of *synchrony of action*, that is, once the mode is activated, all domains or subsystems are involved. Consistent with cognitive theory, the suicidal belief system is evident and potentially amenable to change during periods of activation. Accordingly, activation is critical to treatment progress and success.

The sixth assumption in FVT holds that the severity of the suicidal episode is dependent on the interaction between baseline risk and the severity of the aggravating factors. Again, it is important to recall that baseline risk is essentially the individual's susceptibility to having the suicidal mode triggered, with the suicidal belief system being a critical component in that it is expressed via suicidal thoughts, intent, and related behaviors that prepare for or facilitate a suicidal act. Consistent with previously articulated assumptions about the suicidal mode, the central pathway for suicidality is cognition, the private meaning assigned by the individual for experiences (Rudd et al., 2000). Again, the notion of synchrony of action is important. For example, if the precipitant is a reemergence of depressive symptoms, it is the *interpretation* of this reemergence that is critical (e.g., "It's hopeless; I'll never recover."). Activation of the suicidal mode is dependent on maladaptive meaning constructed and assigned regarding the self, the environmental context, and the future, that is, the suicidal belief system across the four themes referenced earlier. Consistent with this notion of synchrony of action of the suicidal mode, aggravating factors include not only the suicidal belief system, but also the affective, physiological, and behavioral components.

As previously noted, most risk models are categorical in nature, incorporating clinical symptoms of one sort or another. This is one way to understand aggravating factors. What current clinical factors contribute to the patient's risk status? What is the patient thinking? What is the patient feeling emotionally and physiologically? What behavior does the patient mani-

fest? Of importance, the **seventh FVT assumption** states that risk is elevated by aggravating factors for only limited periods of time, such as a few hours, days, or weeks. This is a consequence of the suicidal mode. In short, the body cannot maintain arousal at the highest levels for indefinite periods of time, and, accordingly, risk will naturally resolve to some degree. Even if arousal is diminished to a minimal degree, it may well be enough to move the patient from imminent risk. Again, this process is always accompanied by cognition, that is, the suicidal belief system. For example, limited reductions in arousal might lead to thoughts such as, “Maybe I can tolerate this feeling, I don’t need to kill myself,” countering the core belief themes of helplessness and poor distress tolerance. However, improved symptoms do not always translate to lower risk if the suicidal belief system is active (e.g., “Depression will always be a part of my life, I might as well kill myself”).

In accordance with what was just discussed, the **eighth and final FVT assumption** states that acute risk resolves when aggravating factors are effectively targeted. **In short, you treat the aggravating factors, not baseline (enduring) factors, during periods of crisis and acute risk.** Although this undoubtedly has impact on enduring factors (e.g., cognitive susceptibility), the primary targets are current symptoms (cognitive, affective, and physiological) and behaviors. If those are effectively targeted, then acute risk will resolve. As noted previously, however, risk only returns to baseline level. With reference to the questions previously posed, this is how FVT explains the duration of the suicidal episode. It will last only as long as it takes to treat the aggravating factors effectively. When periods of acute risk endure, it is because the aggravating factors are not effectively targeted. Part of the problem may well be that they are not all identified or recognized by the clinician. As noted before, the suicidal mode has four component parts, and although the suicidal belief system is central, all need to be targeted to resolve a suicidal episode.

CLINICAL IMPLICATIONS:

DIFFERENTIATING ACUTE AND CHRONIC RISK

What are the implications of explaining the process of risk by using fluid vulnerability theory? There are a number of important consequences. First, it should be clear from the assumptions listed in this chapter that **variable baseline risk means that even when some patients recover from a crisis, they can still be at relatively high risk**, with vulnerability that manifests itself across multiple domains. Accordingly, we **need to differentiate between acute and chronic risk for patients.** In other words, all patients have a chronic risk level (i.e., baseline risk). As noted earlier, it is believed that baseline risk for multiple attempters is relatively high.

The implications for clinical practice are considerable. The primary implication is that clinicians need to monitor and record chronic risk factors

(more enduring characterological aspects) in addition to acute risk. For some patients, resolution of an acute suicidal crisis does not mean that suicide risk is low. To the contrary, suicide risk may have only been reduced in marginal fashion. What the clinician needs to do is to make routine the assessment of acute and chronic risk, differentiating the two in all relevant clinical entries. For example, after a patient is discharged from an inpatient unit, risk is certainly not fully resolved after the acute crisis precipitating admission has been targeted. What continues to be of concern are the more enduring and chronic aspects of the individual's case, what I would argue are the component parts of the suicidal mode.

In any event, the clinician needs to note and characterize acute risk carefully at discharge by describing the aggravating factors that were treated across each domain. Similarly, the clinician needs to note that many areas, if not the majority, were not effectively treated in a short inpatient stay. These are the chronic or static factors noted before, comprising the patient's baseline or chronic risk level. For some, such as multiple attempters, this can be fairly high, regardless of effective and efficient treatment. Differentiating chronic and acute risk translates to understanding each suicidal mode component and articulating the patient's vulnerability across each.

A clinical example illustrates best the need to differentiate acute and chronic risk. Let us consider two patients with identical presentations except for a few historical or static factors. For example, say both suffered from a recurrent major depressive problem, were occasionally abusing alcohol, experiencing specific suicidal thoughts with no intent, and had access to a supportive and caring family. One, however, had made four previous and potentially lethal suicide attempts over the past 5 years and the other had no previous suicide attempts. We certainly would be struggling for ways to differentiate their risk status. Although the acute episodes might look very much the same, FVT tells us that after successful treatment of the aggravating factors (across each domain of the suicidal mode), these two patients would return to very different baseline risk levels. The patient with multiple attempts would return to a relatively high level of risk, characterized by a suicidal mode that is easily triggered by a range of stressors both internal and external (across all four domains) and a suicidal belief system likely cutting across all four content themes noted earlier.

For a chronically suicidal patient, it is likely that the suicidal belief system will still be active, even during periods of reduced arousal, emotional upset, and relative behavioral stability. In contrast, the other patient might well return to a baseline risk level that is characterized as minimal or even nonexistent. The patient's suicidal belief system might include a thought such as, "I'll never do that again, life is always worth living" (which could represent modification of an underlying core belief about the value of life, individual worth, and lovability). This is one way of understanding why some patients make a single attempt and fully recover, never to make another

attempt, whereas others go on to make many more. In struggling to differentiate the patient's risk levels, we need to focus on the process of risk for each. How does risk emerge, resolve, and reemerge? Despite the resolution of acute factors for the person with multiple suicide attempts, the chronic problems persist. In many ways, this notion is similar to what Maris (1981) referred to as "suicidal careers."

I recommend that **all clinicians get into the habit of differentiating between acute and chronic risk in every clinical entry for a suicidal patient.** This can be accomplished by simply segmenting the risk assessment section into these two categories. I **also suggest discussing vulnerability for each domain of the suicidal mode.** In addition to a more precise understanding of risk status, it conveys clearly and cogently that risk is a complex construct, one that is not particularly well understood or empirically supported at present. It also makes it clear that there are elements of treatment that are enduring and long term for every suicidal patient, regardless of the pressures of managed care. It can also be argued that such a distinction allows for a more realistic informed consent process, with it incumbent on the clinician to differentiate acute and chronic elements of the patient's suicidality and how those affect the treatment process, content, and duration of care. If our understanding of a patient's suicidality is more precise, it is arguable that our treatment and management activities will be more efficient and effective.

IMPLICATIONS FOR FUTURE RESEARCH

Fluid vulnerability theory has clear implications for all those engaging in research on suicidality. As noted earlier, most (if not all) risk assessment models incorporate risk factors that are derived from empirical research using time frames that are simply inconsistent with risk time frames that are routine in clinical practice. Most clinicians make decisions that involve periods that vary from a few hours to days to weeks. As noted earlier, research on risk factors for suicide uses periods that vary from a year to decades (see Rudd et al., 2000, for review). What is desperately needed are studies that ask and help answer some simple clinical questions. For example, how long does a suicidal state last? How do I know the crisis is over? Are there differences in the duration and severity of suicidal states for multiple attempters?

The fundamental assumptions of FVT will help drive specific hypotheses about risk over acute and chronic time frames. Fluid vulnerability theory should help inform research on clinically relevant risk periods. More specifically, it will help us offer specific hypotheses about how easily a suicidal episode might be triggered, how long it might last, and the likelihood of another episode after recovery. FVT also will help us identify and differentiate high and low risk groups in accordance with the depth and breadth of suicide relevant cognitions (i.e., the suicidal belief system), affective and physiologi-

cal symptoms, and behavioral deficits. At present, this kind of research is in its infancy. We are just now in the process of differentiating those who make multiple attempts from those who make single attempts and those with suicidal ideation (e.g., Rudd et al., 1995). The next step is to advance some of the early work on the parameters of suicidal crises (e.g., Rudd et al., 2000). In particular, we need to track suicidal crises longitudinally, looking at resolution and reactivation over longer periods. Of importance are questions about crisis resolution. Do symptoms remit in clusters or individually? Are some symptoms potentially protective (i.e., the observation of increased risk following remission of neurovegetative symptoms)? What about latent cognitive vulnerability following a suicidal crisis? In other words, do some hopelessness themes endure following crisis resolution? If so, are they significant risk factors? These are but a few questions that surface as important. It is hoped that FVT and the suicidal mode will provide a theoretical foundation that cuts across risk assessment and treatment, providing needed infrastructure for more fruitful scientific study.

REFERENCES

- Alford, B. A., & Beck, A. T. (1997). *The integrative power of cognitive therapy*. New York: Guilford Press.
- Beck, A. T. (1996). Beyond belief: A theory of modes, personality, and psychopathology. In P. Salkovskis (Ed.), *Frontiers of cognitive therapy* (pp. 1–25). New York: Guilford Press.
- Beck, A. T., Emery, G., & Greenberg, R. L. (1985). *Anxiety disorders and phobias: A cognitive perspective*. New York: Guilford Press.
- Clark, D. C., & Fawcett, J. (1992). Review of empirical risk factors for evaluation of the suicidal patient. In B. Bongar (Ed.), *Suicide guidelines for assessment, management, and treatment* (pp. 16–48). New York: Oxford University Press.
- Dowd, E. T., & Courchaine, K. E. (1996). Implicit learning, tacit knowledge, and implications for stasis and change in cognitive psychotherapy. *Journal of Cognitive Psychotherapy*, 10, 163–180.
- Haaga, D. A., Dyck, M. J., & Ernst, D. (1991). Empirical status of cognitive theory of depression. *Psychological Bulletin*, 110, 215–236.
- Hirschfeld, R. M., & Russell, J. M. (1997). Assessment and treatment of suicidal patients. *New England Journal of Medicine*, 337, 910–915.
- Jobes, D. A. (1995). The challenge and promise of clinical suicidology. *Suicide and Life-Threatening Behavior*, 25, 437–449.
- Joiner, T. E., Walker, R. L., Rudd, M. D., & Jobes, D. A. (1999). Scientizing and routinizing the assessment of suicidality in outpatient practice. *Professional Psychology: Research and Practice*, 30, 447–453.
- Litman, R. E. (1991). Predicting and preventing hospital and clinic suicides. *Suicide and Life-Threatening Behavior*, 21, 56–73.

- Maris, R. W. (1981). *Pathways to suicide: A survey of self-destructive behaviors*. Baltimore: Johns Hopkins University Press.
- Rudd, M. D. (2003). Warning signs for suicide? *Suicide and Life-Threatening Behavior*, 33, 99–100.
- Rudd, M. D. (2004). Cognitive therapy for suicidality: An integrative, comprehensive, and practical approach to conceptualization. *Journal of Contemporary Psychotherapy*, 34, 59–72.
- Rudd, M. D., Joiner, T. E., & Rajab, M. H. (1995). Help negation after acute suicidal crisis. *Journal of Consulting and Clinical Psychology*, 6, 499–503.
- Rudd, M. D., Joiner, T. E., & Rajab, M. H. (2000). *Treating suicidal behavior*. New York: Guilford Press.
- Tanney, B. L. (1992). Mental disorders, psychiatric patients, and suicide. In R. W. Maris, A. L. Berman, J. T. Maltzberger, & R. Yuht (Eds.), *Assessment and prediction of suicide* (pp. 277–320). New York: Guilford Press.